

Edmund Weitz

MATHEMATIK

(Skript für den Studiengang *Medieninformatik* an der *HAW Hamburg*)

2. Februar 2026, 14:30

Dieses Skript ist momentan noch „work in progress“ und wird regelmäßig aktualisiert. Sie sollten ab und zu nachschauen, ob es eine neue Version (siehe Datum oben) gibt. Den jeweils letzten Stand finden Sie unter der Adresse <https://weitz.de/files/mi-skript.pdf>.

Falls Sie Fehler entdecken (auch „kleine“ Tippfehler), freue ich mich über eine E-Mail an edmund.weitz@haw-hamburg.de. Für bereits entdeckte Fehler bedanke ich mich bei Marco Ambrosius, Jörg Balzer, Charlotte Benck, Leonard Büttner, Leon Gressler, Alexandra Klan, Meylan Luu, Roger Pearson, Paula Schepler, Emilia Schmitz und Clemens Wagenschein.

Außerdem wird es sicher an manchen Stellen passieren, dass ich Formulierungen verwende, die mir als Mathematiker seit Jahrzehnten geläufig sind, die Ihnen jedoch kryptisch vorkommen. Auch in so einem Fall freue ich mich über eine Mitteilung. Ich möchte gerne, dass alle das Skript verstehen können, bin dabei aber auf Ihre Hilfe angewiesen.

Der gesamte Text ist unter der Creative-Commons-Lizenz [BY-NC-ND 3.0 DE](https://creativecommons.org/licenses/by-nc-nd/3.0/de/) verfügbar. Eine vereinfachte Zusammenfassung dieser Lizenz findet man unter der folgenden Adresse: <https://creativecommons.org/licenses/by-nc-nd/3.0/de/>

Inhaltsverzeichnis

Gebrauchsanweisung	1
I Grundlagen	5
1 Aussagenlogik	7
2 Mengenlehre	18
3 Prädikatenlogik	25
4 „Große“ Mengen	32
5 Arithmetik	39
6 Mengen von Mengen	49
II Zahlentheorie	55
7 Teilbarkeit	55
8 Primzahlen	63
9 Modulare Arithmetik	68
10 Der chinesische Restsatz	79
III Erweiterung der Zahlenwelt	83
11 Ungleichungen und die Zahlengerade	83
12 Die reellen Zahlen	89
13 Runden und Fehlerfortpflanzung	102
14 Die komplexen Zahlen	107
IV Funktionen	119
15 Tupel	119
16 Funktionen	123
V Geometrie	137
17 Elementargeometrie der Ebene	137
18 Die trigonometrischen Funktionen	154
19 Die Arkusfunktionen	164
20 Vektoren	169
VI Lineare Algebra	179
21 Matrizen	179
22 Lineare Gleichungssysteme	186
23 Lineare Abbildungen	197
24 Geometrische Interpretation linearer Abbildungen	201
25 Homogene Koordinaten	206
26 Inverse Matrizen	209
27 Determinanten	212
28 Skalar- und Vektorprodukt	221

29 Raumdrehungen und Quaternionen	230
VII Analysis	237
30 Folgen	237
31 Die Landau-Notation	244
32 Reihen	246
33 Exponentialfunktion und Logarithmus	254
34 Stetige Funktionen	261
35 Differenzieren	265
36 Integrieren	271
VIII Stochastik	281
37 Wahrscheinlichkeit	281
38 Zufallsvariablen	293
39 Hypothesentests	306
Lösungen zu ausgewählten Aufgaben	317
Literatur	411
Index	413
Mathematische Symbole	421

Gebrauchsanweisung

Dass Sie diesen Satz lesen, ist schon mal ein guter Anfang!

Leider lehrt die Erfahrung, dass viele Studierende nicht gerne lesen und Bücher oder Skripte meiden wie der Teufel das Weihwasser. Das ist allerdings keine gute Idee, wenn man erfolgreich studieren will. Menschen, für die Lesen eine Qual ist, rate ich dringend von einem Studium ab!

Vorlesung oder Skript? Beides!

Brauchen Sie überhaupt ein Skript, wenn Sie regelmäßig in der Vorlesung sitzen und mitschreiben? Oder umgekehrt: Müssen Sie noch (eventuell frühmorgens und im Wintersemester womöglich im Dunkeln und bei Minusgraden) in die Vorlesung kommen, wenn Sie das Skript haben? Ich behaupte, dass für fast alle Studierenden die Antwort auf diese beiden Fragen *ja* ist. Vorlesung und Skript decken zwar dieselben Themen ab und theoretisch können Sie mit einem von beiden auskommen. Aber dafür brauchen Sie entweder (nur Skript) sehr viel Selbstdisziplin oder (nur Vorlesung) eine besonders gut ausgeprägte Auffassungsgabe.

Die wöchentliche Vorlesung sorgt für Regelmäßigkeit und gibt damit gleichsam den „Takt“ für das Semester vor. Sie können dem Prof live dabei zusehen, wie man mathematisch arbeitet (und Fehler macht). Sie haben die Gelegenheit, Fragen zu stellen. Sie treffen die anderen Studierenden Ihres Semesters und können sich mit Ihnen austauschen. (Wenn Sie meinen, dass Sie auf all das verzichten können, dann hätten Sie sich vielleicht besser für ein Fernstudium entscheiden sollen, oder?)

Im Skript haben Sie alles noch einmal schwarz auf weiß. Sie können sich die Stellen heraussuchen und nacharbeiten, die Sie nicht gleich verstanden haben. (In der Mathematik versteht so gut wie niemand alles gleich beim ersten Mal!) Sie können Dinge aus den vergangenen Wochen, an die Sie sich nicht mehr so genau erinnern, nachschlagen. Sie können sich auf die kommende Vorlesung vorbereiten, indem Sie ein paar Seiten vorab lesen oder zumindest überfliegen. Und Sie finden evtl. Materialien, die über den in der Vorlesung durchgenommenen Stoff hinausgehen.

Für die Prüfungen gehen wir jedenfalls davon aus, dass Ihnen die Inhalte des Skripts bekannt sind!

Das wichtigste Werkzeug

Wie liest man so ein Skript? Das ist eigentlich schon die falsche Frage. Man *liest* es nicht wie einen Roman im Sessel, sondern man *arbeitet* es durch. Um an einer Mathevorlesung mit Gewinn teilzunehmen, braucht es zwei Zutaten: Verständnis und Routine. Sie können auf Dauer nicht folgen, wenn Sie die Inhalte der Vorwochen nicht verstanden haben, denn in der Regel baut das, was folgt, auf ihnen auf. Deshalb müssen Sie *langsam* lesen, sich dabei quasi selbst ständig über die Schulter schauen und sich immer wieder hinterfragen, ob sie „nur“ gelesen haben oder ob Sie das Gelesene auch mit Ihren eigenen Worten wiedergeben könnten.

Aber selbst dann, wenn Sie alles verstanden haben, brauchen Sie auch Routine. Sie dürfen nicht nur passiv rezipieren, sondern Sie müssen auch aktiv am Stoff arbeiten. Das ist wie mit dem Fahrradfahren oder mit dem Schwimmen: Sie beherrschen es nicht dadurch, dass Sie ein Buch darüber gelesen oder ein Video darüber gesehen haben. Oder wie es der Mathematiker [Carl Runge](#) mal formulierte: „Man lernt das Klavierspielen nicht durch den Besuch von Konzerten.“ Aus diesem Grund sollten Zettel und Bleistift Ihre wichtigsten Werkzeuge werden und beim Durcharbeiten des Skripts immer zur Hand sein. Machen Sie sich Notizen. Fertigen Sie eigenständig Zusammenfassungen der Kapitel an. Gehen Sie Beispiele aus dem Skript selbst durch und glauben Sie nicht einfach, was Sie lesen. Zeichnen Sie eigene Skizzen. Und vor allem: Bearbeiten Sie die Aufgaben! Dafür sind Sie da: die Aufgaben sind ein Service, um Sie beim Lernprozess zu unterstützen. Aber sie helfen Ihnen nur dann, wenn Sie sich mit ihnen beschäftigen.¹

Die Aufgaben

Wenn Sie Abschnitte des Skripts zur Vorbereitung auf die Vorlesung vorab lesen, dann sollten Sie die Aufgaben zu lösen versuchen, *bevor* Sie weiterlesen. Häufig werden Sie die Lösung im folgenden Text finden oder [am Ende des Skripts](#). Aber wenn Sie nicht selbst über die Fragen nachgedacht haben, lernen Sie nichts.

Wenn Sie hingegen das Skript *nach* dem Besuch der Vorlesung noch einmal durcharbeiten, dann kann es sicherlich nicht schaden, sich erneut mit den Aufgaben zu beschäftigen, auch wenn die vielleicht schon besprochen wurden. Versuchen Sie, die Aufgaben zu lösen, ohne Ihre Notizen aus der Vorlesung zur Hilfe zu nehmen.

Unabhängig davon, wie Sie an die Aufgaben herangehen, sollten Sie auf jeden Fall nach getaner Arbeit die Lösungen durchlesen – auch dann, wenn Sie sich sicher sind, dass Ihr Ergebnis korrekt ist. In den Lösungen finden Sie evtl. noch weitere Hinweise oder „Tricks“ bzw. Warnungen.

¹Die Aufgaben sind aber nicht als „Probepfungen“ misszuverstehen. Sie werden in Klausuren mit Sicherheit mit Aufgaben konfrontiert werden, die in diesem Skript nicht zu finden sind, denn Sie sollen ja nicht nur wiederkäuen ...

Beweise

Würde es sich um eine Vorlesung für Studierende der Mathematik an einer Universität handeln, dann müssten alle Aussagen in diesem Skript bewiesen werden. Das ist nämlich gewissermaßen der eigentliche Job von Mathematikerinnen und Mathematikern. Wir machen das hier nicht, weil Sie nicht das Beweisen lernen sollen. Es werden höchstens ab und zu mal Beweise skizziert, wenn mir das didaktisch sinnvoll erscheint. Wenn Sie gerne mehr davon sehen wollen, dann werfen Sie einen Blick in eine ältere Version des Skripts, die man unter der URL <https://weitz.de/files/skript.pdf> findet. Die behandelten Themen sind ungefähr dieselben wie in diesem, aber die Schwerpunkte und die Reihenfolge sind anders. Dafür gibt es zum gesamten Stoff begleitende Videos, die vom Skript aus verlinkt sind.



Grundlagen

Wer sich ernsthaft mit Musik beschäftigen will, wird nicht umhinkommen, die Notenschrift zu erlernen und die Bedeutung von Begriffen wie *Quinte* oder *Triole* zu verstehen. Das ist die Fachsprache, mit der Musikerinnen und Musiker miteinander kommunizieren. Auch die Mathematik hat eine Fachsprache. Dazu gehören die Objekte, die gemeinhin als *Formeln* bezeichnet werden. Darüber hinaus hat die Mathematik auch eine ihr eigene Sprech- und Denkweise. Das wird unser Thema auf den nächsten Seiten sein.

Dieses Kapitel legt das Fundament für alles, was danach folgen wird, und ist deswegen besonders wichtig. Wenn Sie die Sprache des Landes, in dem Sie Ihren Urlaub verbringen, nicht beherrschen, dann können Sie vielleicht mit Händen und Füßen einen Kaffee bestellen, Sie werden sich jedoch nicht mit den Einheimischen unterhalten können. Und wenn Sie die Sprache der Mathematik nicht verstehen, dann sind mathematische Formeln für Sie unverständliche Hieroglyphen, bei deren Anblick Sie regelmäßig verzweifeln.

Man könnte Mathematik auch ohne Formeln betreiben. In früheren Jahrhunderten hat man das gemacht. Aber die Sache wird dadurch nicht einfacher, sondern schwerer, weil natürliche Sprachen wie Deutsch oder Englisch für diesen Zweck zu umständlich sind. Mathematik ohne Formeln führt zu unhandlichen Bandwurm-sätzen, die Missverständnisse geradezu herausfordern.

Aufbau des Skripts bzw. der Vorlesung

Bevor wir loslegen, möchte ich jedoch zunächst erklären, wie Skript und Vorlesung aufgebaut sind. Das Ziel der Mathematik als Wissenschaft ist gewissermaßen die Produktion von verlässlichen Wahrheiten über abstrakte Objekte. Nehmen wir als Beispiel die Behauptung, dass $5^n + 3$ durch 4 teilbar ist, wenn n eine natürliche Zahl ist. Wenn man diese genauer analysiert, erkennt man ein paar Eigenschaften, die so ziemlich alle mathematischen Aussagen gemeinsam haben:

- Es geht nicht um konkrete Dinge wie Magneten, Salzsäure oder Honigbienen. Damit beschäftigen sich Naturwissenschaften wie Physik, Chemie und Bio-

logie. Mathematik ist *keine* Naturwissenschaft und hier geht es um abstrakte Konzepte wie beispielsweise Zahlen. (Zahlen „gibt“ es nicht. Oder sind Sie schon einmal mit einer Sieben zusammengestoßen? Auch einen perfekten [Kreis](#) haben Sie garantiert noch nie gesehen.)

- Sie können sich darauf verlassen, dass die Aussage korrekt ist. Welche Zahl n Sie auch immer nehmen, wenn Sie $5^n + 3$ ausrechnen, wird sich ein Vielfaches von 4 ergeben.
- Das setzt allerdings voraus, dass Einigkeit über die Begriffe besteht. Was genau ist eine natürliche Zahl? Was ist mit 5^n gemeint? Was heißt „durch 4 teilbar“?
- Es geht eigentlich um unendlich viele Aussagen: $5^2 + 3$ ist durch 4 teilbar, $5^{22} + 3$ ist durch 4 teilbar, $5^{1000} + 3$ ist durch 4 teilbar und immer so weiter. Niemand könnte alle diese Fälle durch Nachrechnen überprüfen. Wie kann man sich trotzdem sicher sein, dass das stimmt?

Diese Eigenschaften führen zwangsläufig zu gewissen „Standardzutaten“, denen Sie in jeder Hochschulveranstaltung zur Mathematik – und darum auch in diesem Skript – begegnen werden: Definitionen, Sätze und Beweise.

Definitionen	• In <i>Definitionen</i> wird die Fachsprache festgelegt. Dabei handelt es sich – zumindest im Rahmen der jeweiligen Veranstaltung – um verbindliche Sprachregelungen. Nachdem Begriffe oder Schreibweisen einmal definiert wurden, gilt bei der Verwendung derselben nur noch das, was im Skript steht.
Sätze	• <i>Sätze</i> sind wahre Aussagen über mathematische Objekte. Es gibt außer <i>Satz</i> noch andere Bezeichnungen wie <i>Theorem</i>, <i>Proposition</i>, <i>Lemma</i> oder <i>Korollar</i>, aber die subtilen und eher subjektiven Unterschiede zwischen diesen Begriffen müssen Sie sich nicht merken.¹
Beweise	• <i>Beweise</i> sind die Begründungen dafür, warum Sätze wahr sind – warum man sich beispielsweise darauf verlassen kann, dass $5^n + 3$ immer durch 4 teilbar ist, ohne unendlich viele Divisionen durchzuführen. Im Mathematikstudium sind Beweise die Hauptsache, während wir oft auf Beweise verzichten werden. Sie müssen dann leider einfach glauben, was behauptet wird.²

Auch wenn wir hier nicht das Beweisen lernen, sollte Ihnen für das Verständnis der folgenden Seiten der Unterschied zwischen Definitionen und Sätzen – die sich im Skript auch typographisch voneinander abheben – klar sein. Definitionen können nicht wahr oder falsch sein, weil es sich lediglich um Übereinkünfte handelt. Das prominenteste Beispiel sind die natürlichen Zahlen. Wir werden auf Seite [33](#) definieren, dass die Null eine natürliche Zahl ist. In der Schule war das bei Ihnen vielleicht nicht so und die natürlichen Zahlen fingen bei der Eins an. Und in anderen Büchern wird das evtl. auch anders als hier definiert. Man kann sich

¹Ein Lemma ist quasi ein „kleiner“ Satz, während Korollare „offensichtliche“ Folgerungen sind.

²Wann ein Beweis ausreichend ist, ist übrigens gar nicht so leicht zu beantworten. In der Praxis hängt das vom Kontext und von den Rezipienten ab. In der Theorie beschäftigt sich mit solchen Fragen die [Grundlagenforschung](#).

darüber streiten, was sinnvoller ist, aber keine von beiden Definitionen ist *wahr*. Beide sind lediglich (sich widersprechende) Konventionen.

Sätze sind hingegen immer wahr. Aber sie werden in der Regel in der Fachsprache formuliert und verwenden daher definierte Begriffe. Um ein ganz einfaches Beispiel zu bringen: Wenn a und b natürliche Zahlen sind, dann ist $a + b$ größer als a . In einem Buch, in dem die natürlichen Zahlen bei eins anfangen, ist das ein Theorem: eine wahre Aussage, die man beweisen kann. In diesem Skript wäre es jedoch *kein* Theorem, weil man für b die Zahl Null einsetzen könnte. Trotzdem ist die obige Aussage korrekt, sie müsste nur für dieses Skript anders formuliert werden: Wenn a und b *positive* natürliche Zahlen sind, dann ist $a + b$ größer als a .

In meinem Video *Ist eins eine Primzahl?* gibt es weitere Beispiele zum Unterschied zwischen Definitionen und Sätzen. Zum Glück besteht mit Ausnahme der natürlichen Zahlen bei den meisten häufig verwendeten mathematischen Begriffen Einigkeit. Sie müssen also nicht damit rechnen, dass Sie ständig über Dinge stolpern, die in Buch A richtig, in Buch B jedoch falsch sind. Wenn bei Konzepten Uneinigkeit besteht, werde ich darauf hinweisen.

1. Aussagenlogik

Die Logik als Lehre vom Schlussfolgern, die die Gültigkeit von Argumenten unabhängig von ihren Inhalten untersucht, ist über 2500 Jahre alt und wurde im klassischen Griechenland hauptsächlich durch *Aristoteles* geprägt. Sie galt lange als Teilgebiet der Philosophie. Die moderne mathematische Logik entwickelte sich erst im 19. Jahrhundert unter anderem durch die wegweisenden Arbeiten des Engländers *George Boole*. Die Aussagenlogik, um die es in diesem Abschnitt geht, spielt inzwischen auch eine wichtige Rolle in der Informatik.

Eine **Aussage** (engl. *proposition*) ist ein sprachliches Gebilde, dem man objektiv und zeitlos einen der **Wahrheitswerte** *wahr* („trifft zu“) oder *falsch* („trifft nicht zu“) zuordnen kann. Die beiden Wahrheitswerte schließen sich gegenseitig aus. Eine Aussage ist also immer entweder wahr oder falsch, aber niemals beides.

Definition

Uns werden in der Regel nur Aussagen interessieren, bei denen es um mathematische Sachverhalte geht. Hier ein paar Beispiele für solche Aussagen:

- (i) $6 \cdot 7 = 42$
- (ii) $10 - 2 < 6$
- (iii) Vierzehn ist eine gerade Zahl.
- (iv) Es gibt eine ungerade vollkommene Zahl.

Das erste Beispiel zeigt eine Aussage, die wahr ist, das zweite eine, die falsch ist. Das dritte Beispiel demonstriert, dass mathematische Aussagen keine Formeln sein müssen. (Viele mathematische Fachartikel bestehen größtenteils aus Text und enthalten nur wenige Formeln.) Sie müssen nicht verstehen, worum es im

vierten Beispiel geht. Wichtig ist nur, dass bisher niemand weiß, ob die Aussage wahr oder falsch ist.³ Trotzdem ist es eine Aussage, weil es klare Kriterien für die Beurteilung ihres Wahrheitswertes gibt. (Der Wahrheitswert einer Aussage muss also nicht bekannt sein. Es muss lediglich sinnvoll sein, nach ihrem Wahrheitswert zu fragen.)

Die folgenden Beispiele sind alle *keine* Aussagen.

- (i) $6 \cdot 7$
- (ii) $x^2 - 25 = 0$
- (iii) Die größte bekannte **Primzahl** hat ca. 24 Millionen Dezimalstellen.
- (iv) Geometrie ist interessanter als Algebra.

Das erste Beispiel steht für die Zahl 42. Die ist weder wahr noch falsch. Über den Wahrheitswert des zweiten Beispiels können wir nichts aussagen, weil wir nicht wissen, wofür x steht. Das dritte Beispiel war 2023 noch wahr, seit Ende 2024 ist es jedoch falsch. Der Wahrheitswert ist also nicht – wie oben gefordert – zeitlos. (In der Mathematik beschäftigt man sich nur mit Dingen, die unabhängig von Zeit und Ort sind.) Das vierte Beispiel verletzt die Forderung der Objektivität; manche Menschen werden den Satz für wahr halten, manche für falsch. Auch so etwas spielt in der Mathematik keine Rolle.

Aufgabe 1.1. Welche der folgenden Ausdrücke sind Aussagen?

- (i) $3 \cdot 4 - 12$
- (ii) $3 \cdot 4 = 12$
- (iii) $3 \cdot 5 = 12$
- (iv) $3 \cdot x = 12$
- (v) $3 \cdot 4$ ist eine Primzahl.

Im Folgenden werden wir – wie es auch in der Informatik üblich ist – 1 für *wahr* und 0 für *falsch* schreiben.

Definition

Durch die **Junktoren** $\neg$ („nicht“, **Negation**), $\wedge$ („und“, **Konjunktion**), $\vee$ („oder“, **Disjunktion**), $\Rightarrow$ („wenn, dann“, **Implikation**) und $\Leftrightarrow$ („genau dann, wenn“, **Äquivalenz**) werden Aussagen zu neuen Aussagen verbunden. Die Wahrheitswerte der neuen Aussagen hängen gemäß der folgenden Tabellen ausschließlich von den Wahrheitswerten der verknüpften Aussagen ab und *nicht* davon, worüber diese etwas aussagen.

α und β sind in den Tabellen Platzhalter für beliebige Aussagen. Beachten Sie die Klammern.

³Siehe dazu https://de.wikipedia.org/wiki/Vollkommene_Zahl.

α	$(\neg\alpha)$	α	β	$(\alpha \wedge \beta)$	$(\alpha \vee \beta)$	$(\alpha \Rightarrow \beta)$	$(\alpha \Leftrightarrow \beta)$
0	1	0	0	0	0	1	1
1	0	0	1	0	1	1	0
		1	0	0	1	0	0
		1	1	1	1	1	1

Die Junktoren sind sogenannte *Verknüpfungen*. Damit ist gemeint, dass sie aus Aussagen neue Aussagen machen. In diesem Sinne sind auch $+$ oder $\cdot$ Verknüpfungen, die aus Zahlen wie 2 und 3 neue Zahlen wie $(2 + 3)$ und $(2 \cdot 3)$ machen.

Verknüpfung

Um beispielsweise den Wahrheitswert der Aussage $((7 > 3) \vee (2 \cdot 3 = 5))$ zu ermitteln, stellen wir zunächst fest, dass die Aussage $7 > 3$ wahr ist, während $2 \cdot 3 = 5$ falsch ist. $7 > 3$ wäre in der Tabelle Aussage α und $2 \cdot 3 = 5$ wäre Aussage β . Die dritte Zeile der Tabelle sagt uns nun, dass $((7 > 3) \vee (2 \cdot 3 = 5))$ wahr ist.

Im letzten Absatz wurden der Deutlichkeit halber die beiden Teilaussagen geklammert. Das können wir uns jedoch eigentlich sparen, da keine Gefahr von Verwechslungen besteht: So etwas wie „ $3 \vee 2$ “ würde z. B. keinen Sinn ergeben. Statt $((7 > 3) \vee (2 \cdot 3 = 5))$ können wir also in Zukunft einfach $(7 > 3 \vee 2 \cdot 3 = 5)$ schreiben.

Aufgabe 1.2. Geben Sie die Wahrheitswerte der folgenden sechs Aussagen an:

$$\begin{array}{lll} (\neg(7 > 3)) & (\neg(2 \cdot 3 = 5)) & (2 \cdot 3 = 5 \wedge 7 > 3) \\ (2 \cdot 3 = 5 \Rightarrow 7 > 3) & (7 > 3 \Rightarrow 2 \cdot 3 = 5) & (7 > 3 \Leftrightarrow 2 \cdot 3 = 5) \end{array}$$

Aufgabe 1.3. Warum reicht eine Tabelle mit vier Zeilen aus, um festzulegen, was Junktoren wie $\wedge$ oder $\Rightarrow$ machen?

Eine Reihe von Anmerkungen zu den Junktoren:

- Die Negation $(\neg\alpha)$ hat den gegenteiligen Wahrheitswert von α . Für viele negierte Aussagen gibt es Abkürzungen, die meistens durch das Durchstreichen eines Zeichen realisiert werden. Zum Beispiel ist $3 \neq 4$ eine Abkürzung für $(\neg(3 = 4))$.
- Die Konjunktion $(\alpha \wedge \beta)$ ist nur dann wahr, wenn *beide* Teilaussagen – α und β – wahr sind. In mathematischen Texten wird die Konjunktion häufig abgekürzt geschrieben. Statt

7 ist eine Primzahl und 11 ist eine Primzahl

schreibt man beispielsweise

7 und 11 sind Primzahlen,

wobei hier mit dem Wort *und* natürlich nicht die beiden Zahlen 7 und 11 verknüpft werden, sondern zwei Aussagen.

- Die Disjunktion $(\alpha \vee \beta)$ ist bereits dann wahr, wenn *mindestens* eine der beiden Teilaussagen wahr ist. Anders ausgedrückt ist sie nur dann falsch, wenn *beide* Teilaussagen falsch sind.

- Verwechseln Sie $\vee$ nicht mit dem sogenannten exklusiven Oder⁴, das Sie evtl. aus der Informatik kennen und das im normalen Sprachgebrauch in der Form „entweder oder“ verwendet wird.
- Für die Implikation ($\alpha \Rightarrow \beta$) gibt es u. a. diese Sprechweisen:
 - Wenn α , dann β .
 - Aus α folgt β .
 - α impliziert β .
 - β gilt unter der Voraussetzung α .
 - α ist eine *hinreichende Bedingung* für β .
 - β ist eine *notwendige Bedingung* für α .

hinreichend

notwendig

Sie ist in gewissem Sinne die wichtigste logische Verknüpfung, weil mathematische Aussagen eigentlich immer die Form von Implikationen haben: Eine Aussage β gilt unter der Voraussetzung, dass eine Bedingung α erfüllt sind.

- ($\alpha \Rightarrow \beta$) ist immer wahr, wenn α falsch ist. Ebenfalls ist ($\alpha \Rightarrow \beta$) immer wahr, wenn β wahr ist. Mit anderen Worten: ($\alpha \Rightarrow \beta$) ist nur dann falsch, wenn α wahr und β falsch ist.
- Die Wahrheit von ($\alpha \Rightarrow \beta$) sagt nichts über die zeitliche Abfolge („erst α , dann β “) aus: Bei der vierten Aussage von Aufgabe 1.2 ist $2 \cdot 3 = 5$ nicht „vor $7 > 3$ passiert“. Es geht auch nicht um einen kausalen Zusammenhang („ β , weil α “): In der besagten Aufgabe ist $7 > 3$ nicht deshalb wahr, weil $2 \cdot 3 = 5$ falsch ist. (Siehe dazu auch Aufgabe 1.6.)
- Für die Äquivalenz ($\alpha \Leftrightarrow \beta$) gibt es u. a. diese Formulierungen:
 - α und β sind äquivalent.
 - α genau dann, wenn β .
 - α dann und nur dann, wenn β .
 - α ist *hinreichende und notwendige Bedingung* für β .

Im Englischen sagt man für die dritte Variante „if and only if“ und kürzt das oft als *iff* an. Wenn Sie das in einem Mathebuch sehen, ist das also kein Tippfehler.

- Die Aussage ($\alpha \Leftrightarrow \beta$) ist wahr, wenn α und β denselben Wahrheitswert haben, anderenfalls nicht.

In der Aussagenlogik interessiert man sich nur für die logische Struktur von Aussagen und nicht für deren Inhalte. Beispielsweise kann man zeigen, dass eine Aussage der Form $((\alpha \wedge \beta) \Rightarrow \alpha)$ immer wahr ist – unabhängig davon, was α und β aussagen. Darum entwickeln wir nun einen entsprechenden Kalkül, mit dessen Hilfe man die Wahrheitswerte zusammengesetzter Aussagen ermitteln kann.⁵

⁴Siehe <https://de.wikipedia.org/wiki/Kontravalenz>.

⁵Mit dem Wort *Kalkül* bezeichnet man ein System von Regeln. Es kommt vom lateinischen Wort *calculus*, das für einen Spielstein steht. Es geht also quasi um „Spielregeln“.

Buchstaben wie A, B, C und so weiter (oder auch $x, y, z, \dots$ oder $p_1, p_2, p_3, \dots$) werden **Aussagenvariablen** genannt. Jede Aussagenvariable ist eine (**atomare**) **Formel der Aussagenlogik**. Sind α und β Formeln der Aussagenlogik, dann sind auch $(\neg\alpha)$, $(\alpha \wedge \beta)$, $(\alpha \vee \beta)$, $(\alpha \Rightarrow \beta)$ und $(\alpha \Leftrightarrow \beta)$ **Formeln der Aussagenlogik**. Nur Zeichenketten, die durch endlich viele Anwendungen dieser Regeln entstanden sind, sind Formeln (der Aussagenlogik).

Definition

Den Zusatz „der Aussagenlogik“ werden wir im Folgenden weglassen.

Wie ist das gemeint? Nach Definition sind beispielsweise A, C und E Formeln. Mit einmaliger Anwendung der Regeln (auf $\alpha = C$ und $\beta = E$) ist auch $(C \Rightarrow E)$ eine Formel. Analog sind $(\neg A)$ und $(C \wedge C)$ Formeln. Mit erneuter Anwendung (auf $\alpha = (C \Rightarrow E)$ und $\beta = A$) erhält man dann z. B. die Formel $((C \Rightarrow E) \wedge A)$. So würde man auch $((\neg A) \Leftrightarrow (C \wedge C))$ erhalten.

Die Aussagenvariablen stehen stellvertretend für Aussagen wie $3 < 7$, von denen uns in diesem Zusammenhang nicht interessiert, worüber sie etwas aussagen, sondern lediglich, ob sie wahr oder falsch sind.

Aufgabe 1.4. Welche der folgenden acht Zeichenketten sind Formeln? Begründen Sie jeweils Ihre Antwort.

$$\begin{array}{cccc} (E \wedge \wedge E) & (EE) & (A \wedge B \wedge C) & E \\ (A \wedge (B \wedge C)) & ((A \wedge C)) & A \wedge B & (A + B) \end{array}$$

Nebenbei haben Sie hier ein erstes Beispiel für eine sogenannte *formale Sprache* kennengelernt: Es gibt einen Vorrat von Zeichen (in diesen Fall die Großbuchstaben, die Klammern und die Junktoren) sowie Regeln, nach denen aus diesen Zeichen „erlaubte“ Zeichenketten gebildet werden können. So etwas spielt in der *theoretischen Informatik* eine wichtige Rolle.

formale Sprache

Eine **Belegung** für eine Ansammlung von Aussagenvariablen ist eine Zuordnung von Wahrheitswerten zu diesen Variablen.

Definition

Zum Beispiel erhält man eine Belegung der Variablen A, B und D , indem man B den Wert 1 und A und D den Wert 0 zuordnet.

Aufgabe 1.5. Wie viele verschiedene Belegungen gibt es für drei Variablen? Wie viele gibt es für n Variablen?

Jeder Formel und jeder Belegung der in ihr vorkommenden Aussagenvariablen wird nun sinnvollerweise ein Wahrheitswert zugeordnet, indem die Formel „von innen nach außen“ gemäß der Tabelle zur Definition von Seite 8 ausgewertet wird. Nehmen wir beispielsweise die Formel $((B \vee C) \wedge A)$ und die Belegung, die C den Wert 1 und A und B den Wert 0 zuordnet, so erhält die Teilformel $(B \vee C)$ den Wert 1 aus der Tabellenspalte für die Disjunktion (zweite Zeile) und die gesamte

Formel den Wert 0 aus der Spalte für die Konjunktion (dritte Zeile). Das kann man tabellarisch so darstellen:

A	B	C	$(B \vee C)$	$((B \vee C) \wedge A)$
0	0	1	1	0

Mit etwas Übung spart man sich Schreibearbeit, wenn man die Formel nur einmal hinschreibt und die Teilauswertungen direkt unter die Junktoren schreibt. Man muss dabei nur aufpassen, dass man in der richtigen Reihenfolge vorgeht und nicht den Überblick verliert.

A	B	C	$((B \vee C) \wedge A)$
0	0	1	1 0

Wahrheitstafel

Führt man diesen Vorgang für alle möglichen Belegungen der in einer Formel vorkommenden Variablen durch, so ergibt das die *Wahrheitstafel* für diese Formel. Hier ein Beispiel für die Formel $((A \vee B) \Rightarrow A)$:

A	B	$((A \vee B) \Rightarrow A)$
0	0	1
0	1	0
1	0	1
1	1	1

Definition

Eine Formel wird **erfüllbar** genannt, wenn es mindestens eine Belegung der in ihr vorkommenden Variablen gibt, die ihr den Wahrheitswert 1 zuordnet. Erhält die Formel den Wahrheitswert 1 für *jede* Belegung, dann spricht man von einer **Tautologie**. Zwei Formeln werden als **äquivalent** bezeichnet, wenn sie für alle Belegungen der in ihnen vorkommenden Variablen dieselben Wahrheitswerte erhalten.

Zum Beispiel zeigt die Wahrheitstafel vor dieser Definition, dass $((A \vee B) \Rightarrow A)$ erfüllbar, aber keine Tautologie ist.

Aufgabe 1.6. Überzeugen Sie sich durch Wahrheitstafeln, dass $(A \Rightarrow B)$ äquivalent zu $((\neg A) \vee B)$ und zu $(\neg(A \wedge (\neg B)))$ ist.

Aufgabe 1.7. Überzeugen Sie sich durch Wahrheitstafeln, dass die folgenden Paare von Formeln jeweils äquivalent sind:

- (i) A und $(\neg(\neg A))$
- (ii) $(A \wedge B)$ und $(B \wedge A)$
- (iii) $(A \vee B)$ und $(B \vee A)$
- (iv) $(A \Leftrightarrow B)$ und $(B \Leftrightarrow A)$
- (v) $((A \wedge B) \wedge C)$ und $(A \wedge (B \wedge C))$
- (vi) $((A \vee B) \vee C)$ und $(A \vee (B \vee C))$
- (vii) $((A \Rightarrow B) \wedge (B \Rightarrow A))$ und $(A \Leftrightarrow B)$

In der Definition der Formeln der Aussagenlogik wurden vorsichtshalber in allen Regeln Klammern verwendet, damit später die Reihenfolge der Auswertung einer Formel („von innen nach außen“) eindeutig definiert war. Entsinnen Sie sich, dass es z. B. nicht egal ist, ob man $5 + (3 \cdot 4)$ oder $(5 + 3) \cdot 4$ ausrechnet. Bei den Grundrechenarten gibt es jedoch Regeln wie „Punkt- vor Strichrechnung“ zur Klammerersparnis. Solche Regeln werden wir nun auch einführen:

- Das äußerste Klammernpaar kann bei Formeln grundsätzlich weggelassen werden. Statt $(A \Rightarrow (B \Rightarrow C))$ schreiben wir z. B. in Zukunft $A \Rightarrow (B \Rightarrow C)$.
- Die Junktoren, die in der Tabelle zur Definition von Seite 8 weiter links stehen, haben höhere *Priorität*. $\wedge$ hat also beispielsweise höhere Priorität als $\Leftrightarrow$ und $\neg$ hat höhere Priorität als $\vee$. Bei fehlenden Klammern „binden“ die Junktoren mit höherer Priorität stärker. Beispielsweise ist $A \wedge B \Leftrightarrow C$ eine Abkürzung für $(A \wedge B) \Leftrightarrow C$ und $\neg A \Rightarrow B$ eine für $(\neg A) \Rightarrow B$.
- Nach Aufgabe 1.7 kann man beruhigt $A \wedge B \wedge C$ schreiben, weil $((A \wedge B) \wedge C)$ und $(A \wedge (B \wedge C))$ zwar nicht dieselbe Formel, aber zumindest äquivalent sind. Das gilt ebenso für $\vee$ und $\Leftrightarrow$, aber *nicht* für $\Rightarrow$! (Siehe dazu Aufgabe 1.11.) Wenn man formal korrekt vorgehen will, so setzt man fest, dass in solchen Fällen Linksklammerung gemeint ist: $A \wedge B \wedge C$ ist also eine Abkürzung $(A \wedge B) \wedge C$.
- In vielen Büchern wird eine Formel der Form $\alpha \Rightarrow \beta \Rightarrow \gamma$ als Abkürzung für $\alpha \Rightarrow (\beta \Rightarrow \gamma)$ betrachtet. Wir werden in solchen Fällen (und auch bei $\Leftrightarrow$) vorsichtshalber immer Klammern setzen.
- Wie beim Rechnen gilt auch hier: Verwenden Sie im Zweifelsfall lieber ein „überflüssiges“ Klammernpaar statt unnötige Fehler zu machen.

Priorität

Die Regeln zur Klammerersparnis werden wir in Zukunft nicht nur für Formeln, sondern für beliebige Aussagen verwenden; also auch dann, wenn mit Junktoren nicht Aussagenvariablen, sondern Aussagen wie $3 < 7$ verknüpft werden.

Aufgabe 1.8. Setzen Sie in der Formel

$$A \vee B \Rightarrow \neg A \vee C \wedge B \Leftrightarrow A \vee C$$

alle Klammern, die nach den obigen Regeln weggelassen wurden.

Aufgabe 1.9. Entfernen Sie in der Formel

$$(((C \wedge B) \vee (A \wedge B)) \Rightarrow ((A \vee B) \wedge (B \vee C)))$$

so viele Klammern wie möglich.

Aufgabe 1.10. In Aufgabe 1.7 wird nicht behauptet, dass $A \Rightarrow B$ und $B \Rightarrow A$ äquivalent sind. Warum nicht?

Aufgabe 1.11. Überzeugen Sie sich, dass $A \Rightarrow (B \Rightarrow C)$ und $(A \Rightarrow B) \Rightarrow C$ *nicht* äquivalent sind.⁶

⁶Das ist wie beim Rechnen, wo $5 - (3 - 1)$ und $(5 - 3) - 1$ zu unterschiedlichen Ergebnissen führen.

Aufgabe 1.12. Überzeugen Sie sich, dass $A \Leftrightarrow (B \Leftrightarrow C)$ und $(A \Leftrightarrow B) \Leftrightarrow C$ äquivalent sind. Vergleichen Sie mit Aufgabe 1.11.

Aufgabe 1.13. Stellen Sie sicher, dass Sie die folgende Aussage verstanden haben und dass Ihnen auch klar ist, *warum* sie korrekt ist: Zwei Formeln α und β sind genau dann äquivalent, wenn $\alpha \Leftrightarrow \beta$ eine Tautologie ist.

Beachten Sie, dass in Aufgabe 1.13 die Aussagenlogik quasi auf zwei Ebenen vorkommt. Inhaltlich geht es um die Äquivalenz von Formeln, gleichzeitig sind die beiden Aussagen über α und β aber selbst auch äquivalent, denn es geht ja darum, dass die eine *genau dann* wahr ist, wenn die andere wahr ist.

Von der Korrektheit der folgenden Aussage kann man sich ebenfalls durch Wahrheitstafeln überzeugen (und das sollten Sie zur Übung auch tun). Benannt ist sie nach dem englischen Mathematiker **Augustus De Morgan** aus dem 19. Jahrhundert.

Lemma 1.1

De-morgansche Gesetze

Die Formeln $\neg(A \wedge B)$ und $\neg A \vee \neg B$ sind äquivalent.

Die Formeln $\neg(A \vee B)$ und $\neg A \wedge \neg B$ sind ebenfalls äquivalent.

Man kann sich diese Regeln vielleicht dadurch merken, dass die Negation, wenn man sie „nach innen“ zieht, Konjunktion und Disjunktion vertauscht.

Aufgabe 1.14. Begründen Sie, dass die Formel $(A \Rightarrow B) \Leftrightarrow (B \Rightarrow A)$ erfüllbar, aber keine Tautologie ist.

Aufgabe 1.15. In Aufgabe 1.6 hat man gesehen, dass man $\Rightarrow$ in gewissem Sinne gar nicht braucht, weil man dasselbe auch mithilfe von $\wedge$ und $\neg$ darstellen kann. Geben Sie entsprechend Formeln an, die äquivalent zu $A \vee B$ bzw. $A \Leftrightarrow B$ sind und in denen nur $\wedge$ und $\neg$ vorkommen. Überprüfen Sie Ihre Antworten mithilfe von Wahrheitstafeln.

Aufgabe 1.16. Überprüfen Sie mit einer Wahrheitstafel, ob

$$(A \vee B \Rightarrow C) \Leftrightarrow (A \Rightarrow C) \wedge (B \Rightarrow C)$$

erfüllbar oder sogar eine Tautologie ist.

Aufgabe 1.17. Wir betrachten die Formel $(\neg A \vee B) \wedge P \Rightarrow (A \Leftrightarrow B)$. Wie kann man aus dieser eine Tautologie machen?

- (i) Indem man P durch $(B \Rightarrow A)$ ersetzt.
- (ii) Indem man P durch $(\neg A \wedge B)$ ersetzt.
- (iii) Indem man P durch $\neg(B \wedge \neg A)$ ersetzt.

Aufgabe 1.18. Wenn Sie noch Übung im Umgang mit Wahrheitstafeln brauchen, dann überzeugen Sie sich, dass $F \wedge (F \vee G) \Leftrightarrow F$ und $F \vee (G \wedge H) \Leftrightarrow (F \vee G) \wedge (F \vee H)$ Tautologien sind und dass das auch dann noch gilt, wenn man in den beiden Formeln die Rollen von $\wedge$ und $\vee$ vertauscht.

Zum Üben folgen nun Aufgaben, in denen es nicht um mathematische Aussagen, sondern um andere Dinge geht. Es handelt sich dabei um sogenannte **Logikrätsel**. Entscheidend ist dabei immer die Voraussetzung, dass – anders als evtl. in der Realität – jede Aussage eindeutig wahr oder falsch ist und Begriffe wie „oder“ und „wenn, dann“ so wie in diesem Abschnitt gemeint sind.

Aufgabe 1.19. Ein altes Sprichwort lautet: „Wenn der Hahn kräht auf dem Mist, dann ändert sich das Wetter oder es bleibt, wie es ist.“ Wie sieht das als Formel der Aussagenlogik aus und was kann man über den Wahrheitsgehalt dieses Sprichwortes sagen?

Aufgabe 1.20. Vor Ihnen stehen zwei verschlossene Kisten. Sie wissen, dass sich in jeder der Kisten genau ein Stück Obst befindet, und zwar jeweils ein Apfel oder eine Orange. Sie wissen aber nicht, welche Sorte Obst sich in welcher Kiste befindet, und es kann auch sein, dass in beiden Kisten ein Apfel oder in beiden eine Orange ist. Ferner klebt auf jeder Kiste ein Schild. Auf dem Schild auf der ersten Kiste steht: „In dieser Kiste ist eine Orange und in der anderen ist ein Apfel.“ Auf dem Schild auf der zweiten Kiste steht: „Es befindet sich nicht in beiden Kisten dieselbe Sorte Obst.“ Schließlich wird Ihnen noch mitgeteilt, dass eine der beiden Aussagen wahr und die andere falsch ist, aber Ihnen wird natürlich nicht gesagt, welche der beiden Aussagen die wahre ist. Sie lieben Orangen. Welche Kiste sollten Sie öffnen, wenn Sie nur eine öffnen dürfen? (Oder können Sie vielleicht schließen, dass in beiden Kisten Äpfel sind?)

Aufgabe 1.21. Haben Sie in Aufgabe 1.20 das Schild auf der zweiten Kiste anders übersetzt? Wie könnte man es noch machen?

Aufgabe 1.22. Zwei der folgenden vier Sätze machen (abgesehen von grammatikalischen Feinheiten) rein logisch exakt dieselbe Aussage. Welche beiden sind es?

- (i) Wenn der HSV die Champions League gewinnt, dann springe ich in die Alster.
- (ii) Wenn der HSV nicht die Champions League gewinnt, dann springe ich nicht in die Alster.
- (iii) Wenn ich nicht in die Alster springe, dann hat der HSV die Champions League nicht gewonnen.
- (iv) Wenn der HSV die Champions League gewinnt, dann springe ich nicht in die Alster.

Aufgabe 1.23. Sie haben eine Druckerei beauftragt, Spielkarten nach den folgenden Regeln zu bedrucken:

- (i) Auf einer Seite steht ein Buchstabe, auf der anderen eine natürliche Zahl.
- (ii) Wenn auf der Buchstabenseite ein Vokal steht, dann muss auf der Zahlenseite eine gerade Zahl stehen.

Vor Ihnen liegen vier Probedrucke, auf denen Sie die Zeichen U, 4, P und 7 sehen. Sie sind sich sicher, dass die Druckerei die erste Regel bei allen vier Karten befolgt hat, aber Sie wollen überprüfen, ob auch die zweite Regel immer eingehalten wurde. Wie viele und welche der vier Karten müssen Sie dafür umdrehen? (Sie sind faul und wollen möglichst wenige Karten umdrehen!)

Aufgabe 1.24. Kommissar Gödel muss einen Mordfall lösen und hat drei Tatverdächtige: Gauß, Euler und Hilbert. Aufgrund der bisherigen Ermittlungen weiß er die folgende Dinge mit Sicherheit:

- (i) Es kann einen, aber auch mehrere Täter gegeben haben.
- (ii) Ist Hilbert schuldig, so war auch Euler einer der Täter.
- (iii) Stellen sich Gauß oder Hilbert als Täter heraus, dann ist Euler unschuldig.
- (iv) Sind aber Euler oder Hilbert unschuldig, dann muss Gauß an der Tat beteiligt gewesen sein.

Helfen Sie dem Kommissar mit Methoden der Aussagenlogik.

Aufgabe 1.25. Eine Astronautin besucht einen Planeten, dessen Bewohner – die *Froods* genannt werden – entweder blau oder grün sind. Die blauen Froods sagen immer die Wahrheit, die grünen Froods lügen immer. Die Astronautin trifft die beiden Froods Andrew und Bert, aber es ist zu dunkel, um ihre Farben zu erkennen. Sie fragt Andrew, ob er und Bert beide blau seien. Andrew antwortet, aber die Astronautin weiß danach noch nicht, welche Farbe die beiden haben. Darauf fragt sie Andrew, ob er und Bert dieselbe Farbe haben. Nachdem Andrew auch diese Frage beantwortet hat, kennt sie die Farben der beiden. Sie auch?

Bemerkungen

In der Mathematik ist es üblich, Buchstaben als Platzhalter für mathematische Objekte zu verwenden. Das macht man bereits in der Schule und man bezeichnet diese Buchstaben als *Variablen* oder *Veränderliche*. Meistens werden Buchstaben des *lateinischen Alphabets* (a bis z bzw. A bis Z) verwendet, manchmal aber auch *griechische* wie α , β , γ und so weiter. Dabei ist zu beachten, dass Variablen nicht nur für Zahlen stehen können, sondern auch für andere Dinge. Schon in der Schule werden beispielsweise auch für *Funktionen* oder Winkel Buchstaben verwendet. Im vorherigen Abschnitt standen Variablen wie α und β für Formeln der Aussagenlogik.

Es gibt gewisse Konventionen (z. B. m und n für *natürliche Zahlen*, x und y für *reelle Zahlen*, f und g für Funktionen, Pfeile wie $\vec{v}$ oder Fettdruck wie $\mathbf{v}$ für *Vektoren*, aber diese werden nicht durchgehend eingehalten, weil es erstens keine festen Regeln und zweitens schlicht und einfach zu wenige Buchstaben gibt.

Mathematische Variablen werden im Druck in der Regel in einer speziellen *Kursivschrift* gesetzt, um die Lesbarkeit zu erhöhen. Man beachte zum Beispiel den Unterschied zwischen x (falsch) und x (richtig). Manchmal ist es auch praktisch, *indizierte* Buchstaben wie a_1 , a_2 , a_3 und so weiter zu verwenden, weil man dann nicht durch die Größe des Alphabets eingeschränkt ist, sondern so viele Symbole zur Verfügung hat, wie man braucht. Das ist selbstverständlich so gemeint, dass es sich um lauter verschiedene Variablen handelt.

Kommt ein und dieselbe Variable mehrfach vor, so steht sie natürlich immer für *dasselbe* Objekt. Beispielsweise stimmt die Formel $a + 1 = 1 + a$ offensichtlich nicht mehr, wenn man links vom Gleichheitszeichen 3 für a einsetzt, rechts jedoch 4. Umgekehrt ist es jedoch ein häufig auftretender Denkfehler, dass unterschiedliche Variablen für verschiedene Dinge stehen. Sie *können* für verschiedene Dinge stehen, *müssen* das aber nicht. Die *binomische Formel* $(a + b)^2 = a^2 + 2ab + b^2$ stimmt z. B. auch noch dann, wenn man für a und für b den Wert 7 einsetzt.

Ein weiterer wichtiger Punkt, der Missverständnisse verhindern kann, ist die Betonung der Zeitlosigkeit in der Definition am Anfang des Abschnitts. Es ist sinnvoll, sich die gesamte Welt der Mathematik als statisch vorzustellen: Eine Formel wie $2 + 5 = 7$ ist nicht im Sinne eines Prozesses („ $2 + 5$ ergibt 7“) zu verstehen, wie es Kinder manchmal tun, die deswegen den Ausdruck $7 = 2 + 5$ verwirrend finden. Die bessere Analogie ist, dass $2 + 5$ und 7 zwei unterschiedliche Darstellung desselben Objekts sind und dass das Gleichheitszeichen diese Identität, die „schon immer“ galt, zum Ausdruck bringt.

Wir haben in Aufgabe 1.10 gesehen, dass $A \Rightarrow B$ und $B \Rightarrow A$ *nicht* äquivalent sind. So eine Äquivalenz wird jedoch von vielen Menschen implizit angenommen und führt zu Fehlschlüssen. Darum sei hier noch einmal deutlich darauf hingewiesen, dass man daraus, dass B aus A folgt, *nicht* schließen kann, dass A aus B folgt. Ein einfaches Beispiel: Wenn eine Zahl durch 4 teilbar ist, dann ist sie gerade. Das stimmt. Aber das bedeutet nicht, dass jede gerade Zahl durch vier teilbar ist. Das müsste dann ja auch für 2, 6, 10 und so weiter gelten.

Es gibt diverse Alternativen zu den Symbolen, die in diesem Skript für Junktoren verwendet werden, z. B. $\bar{A}$ statt $\neg A$ oder $A \cdot B$ statt $A \wedge B$, aber es würde wohl nur Verwirrung stiften, die alle aufzuführen. In vielen weitverbreiteten Programmiersprachen, deren Syntax von der von C abgeleitet ist, verwendet man für $\neg$, $\wedge$ und $\vee$ die Zeichen `!`, `&&` und `||`.

Bitte beachten Sie, dass die Symbole für Junktoren zwar in *Formeln* Sinn ergeben, dass sie jedoch *nicht* als Abkürzungen für Wörter gedacht sind. Es wäre furchtbar schlechter Stil, so etwas wie „ $p \wedge q$ sind positiv“ oder „ x ist $\neg$ rational“ zu schreiben; jede Fachzeitschrift würde das ablehnen.

Wie Ihnen schon aufgefallen sein dürfte, sagt man auch, dass eine Aussage *gilt*, wenn man ausdrücken will, dass sie wahr ist. Es ist also beispielsweise üblich zu sagen, dass $6 \cdot 7 = 42$ gilt.

★ Logik im Computer

Abschnitte oder Aufgaben, die *mit einem Stern* gekennzeichnet sind, sind optional. Wenn Sie Angst haben, aus Versehen etwas zu lernen, das Sie für die Prüfung nicht brauchen, sollten Sie sie überspringen. Andererseits steht so etwas nicht deshalb im Skript, weil ich es aus Langeweile eingefügt habe, sondern weil ich glaube, dass es interessant ist und Ihnen eventuell beim Verstehen und ganz allgemein im Studium helfen wird. ★

Ich werde zwischendurch ab und zu zeigen, wie die Mathematik, um die es in diesem Skript geht, auf dem Computer eingesetzt werden kann. Dazu verwende ich die Programmiersprache PYTHON.⁷

Wenn Sie in PYTHON

⁷Das Skript kann Ihnen allerdings keine umfassende Einführung in PYTHON liefern. Es gibt dafür z. B. in der [Bibliothek der HAW Hamburg](#) genügend kostenlose Materialien. Unter anderem können Sie PYTHON auch mit meinem Buch [Konkrete Mathematik \(nicht nur\) für Informatiker](#) lernen.

```
42 > 23
```

eingeben, dann erhalten Sie die Antwort `True`. Das ist in dieser Programmiersprache eine **Konstante**, die für *wahr* steht. Gleichzeitig erkennt man an der Antwort, dass `42 > 23` als Aussage aufgefasst und ihr ein Wahrheitswert zugeordnet wurde. Die Eingabe

```
not 42 > 23
```

liefert die Antwort `False`. Darin sieht man nicht nur, dass `False` die Konstante mit der Bedeutung *falsch* ist und dass `not` für die Negation steht, sondern man lernt auch gleich etwas über die **Priorität der Operatoren**. (Nämlich was?)

Die Eingabe `type(True)` führt zur Ausgabe `bool`.⁸ Das ist der **Datentyp**, den `True` und `False` haben – und außer ihnen kein anderer Wert.

Schließlich noch ein einfaches Beispiel dafür, wie Konjunktion und Disjunktion in PYTHON aussehen. Probieren Sie es aus!

```
for x in [10,30,50]:
    if x > 23 and x < 42:
        print(f"{x} > 23 UND {x} < 42")
    if x > 23 or x < 42:
        print(f"{x} > 23 ODER {x} < 42")
```

2. Mengenlehre

Mengen sind heutzutage die Grundbausteine der gesamten Mathematik, obwohl sie im Vergleich zum Alter dieser Wissenschaft junge Hüpfen sind. Insbesondere ist die Mengenschreibweise ein wesentlicher Teil der mathematischen Fachsprache. Wenn man sie nicht beherrscht, wird man große Teile der Fachliteratur, auch in der Informatik, nicht verstehen können.⁹ Zudem ist die Mengenlehre ein wichtiger Bestandteil der **theoretischen Informatik** und Mengen können beim Programmieren eine nützliche Datenstruktur sein.

Das mathematische Teilgebiet der **Mengenlehre** wurde erst nach 1870 von dem Deutschen **Georg Cantor** begründet, entwickelte sich dann aber ziemlich stürmisch. Der faszinierendste Aspekt der Mengenlehre sind die von Cantor entdeckten **unterschiedlichen Unendlichkeiten**. Mit denen werden wir uns allerdings erst später und nur kurz beschäftigen.

Definition

Eine **Menge** (engl. *set*) ist eine (ungeordnete) Ansammlung von (mathematischen) Objekten, die man die **Elemente** der Menge nennt. Ist ein Objekt *a* Element einer Menge *M*, so schreibt man $a \in M$.

⁸Benannt nach dem bereits erwähnten **George Boole**.

⁹Das gilt auch für dieses Skript!

Statt „ a ist ein Element von M “ sind auch Formulierungen wie „ a ist in M enthalten“, „ a gehört zu M “, „ M enthält a “ und so weiter üblich. Wie auf Seite 9 bereits erwähnt verwendet man die durchgestrichene Form $a \notin M$ als Abkürzung für die Negation $\neg(a \in M)$.

Intuitiv kann man sich eine Menge vielleicht wie einen Behälter (etwa eine Tasche oder einen Karton) vorstellen, bei dem uns nur der Inhalt (die Elemente) interessiert. In den ersten Beispielen werden die Objekte, von denen in der Definition die Rede ist, Zahlen sein. Prinzipiell kann jedoch *jedes* mathematische Objekt Element einer Menge sein; dazu gehören beispielsweise Vektoren, Kreise, Matrizen und Funktionen und auch – wie wir noch sehen werden – Mengen selbst.

Überschaubare Mengen kann man dadurch notieren, dass man alle ihre Elemente durch Kommata getrennt zwischen geschwungenen Klammern (sogenannten *Mengenklammern*) aufführt. Das ist die *aufzählende Mengenschreibweise*, die so aussieht:

$$A = \{23, 42, \pi\}.$$

Für dieses Beispiel gilt unter anderem $23 \in A$ und $13 \notin A$. Und nicht nur für 13, sondern für *jedes* mathematische Objekt x außer 23, 42 und π gilt $x \notin A$.

Eine Menge wie $\{12\}$, die nur ein einziges Element enthält, bezeichnet man als *Singleton*. Ganz wesentlich für das Verständnis ist, dass die Menge $\{x\}$ und ihr Element x *verschiedene* Objekte sind: $x \neq \{x\}$.

Allgemein ist $x \in M$ für jede Menge M und jedes mathematische Objekt x eine Aussage im Sinne der Definition auf Seite 7 (die dementsprechend wahr oder falsch sein kann). Die definierende Eigenschaft von Mengen ist, dass das die *einzigste* „Frage“ ist, die Mengen „beantworten“ können. Eine Menge M „weiß“ für jedes Objekt x , ob x ein Element von M ist, aber ansonsten „weiß“ M nichts!

Das hat zwei wichtige Konsequenzen:

- (i) Mengen kennen keine Reihenfolge. $A = \{23, 42, \pi\}$ und $B = \{\pi, 23, 42\}$ sind beispielsweise dieselbe Menge: $A = B$. (Denn A und B geben immer dieselben „Antworten“, wenn man sie fragt, ob etwas ein Element von ihnen ist.)
- (ii) Ein Objekt kann nur einmal Element einer Menge sein. $C = \{23, 23, 42, \pi\}$ kann man zwar schreiben, aber C ist immer noch dieselbe Menge wie A : $A = C$. (Denn A und C geben immer dieselben „Antworten“, wenn man sie fragt, ob etwas ein Element von ihnen ist.)

Es ist nicht verboten, dass eine Menge gar kein Element enthält. Das würde dann $\{\}$ geschrieben werden, wofür allerdings das Zeichen $\emptyset$ gebräuchlicher ist. Da wir gerade gelernt haben, dass Mengen durch die Angabe ihrer Elemente eindeutig bestimmt sind, gibt es nur eine solche Menge. Sie wird als *die leere Menge* bezeichnet.

aufzählende
Mengenschreibweise

Singleton

$\emptyset$
leere Menge

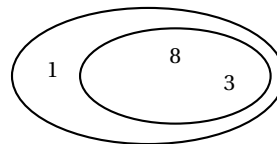
Sind A und B Mengen, so wird A **Teilmenge** (engl. *subset*) von B genannt, wenn jedes Element von A auch Element von B ist. Man schreibt dafür $A \subseteq B$.

Definition

Alternativ wird dann B auch als **Obermenge** von A bezeichnet und dafür $B \supseteq A$ geschrieben.

Zum Beispiel ist $A = \{3, 8\}$ eine Teilmenge von $B = \{1, 3, 8\}$: A hat die beiden Elemente 3 und 8, und beide sind auch Elemente von B . B ist jedoch *nicht* Teilmenge von A , denn das Element 1 von B ist kein Element von A .

Häufig kann man sich die Beziehungen zwischen Mengen grafisch durch sogenannte *Mengendiagramme* klarmachen. Dafür werden die einzelnen Mengen als Kreise, Ellipsen oder ähnliche geometrische Gebilde dargestellt und man stellt sich vor, dass die Elemente der Mengen sich im Inneren dieser Figuren befinden. In der folgenden Skizze steht die große Ellipse für die Menge B aus dem gerade besprochenen Beispiel und die kleine für die Menge A .



Obwohl auch die Elemente dargestellt wurden, sind für das Verständnis der Zusammenhänge eigentlich nur die umrandeten Flächen relevant.

Aufgabe 2.1. Gegeben seien die Mengen $X = \{10, 20, 33\}$, $Y = \{10, 20\}$ und $Z = \{20, 33\}$. Es gibt insgesamt neun Möglichkeiten, eine Teilmengenbeziehung auszudrücken: $X \subseteq X$, $X \subseteq Y$, $Y \subseteq X$, $X \subseteq Z$ und so weiter. Geben Sie für jede von diesen neun Aussagen an, ob sie wahr oder falsch ist.

Dass der folgende Satz stimmt, folgt direkt aus der Definition der Teilmengenbeziehung. Machen Sie sich jeweils – ggf. anhand von Beispielen – klar, dass Sie die einzelnen Aussagen verstanden haben und dass Ihnen klar ist, warum sie wahr sind.

Lemma 2.1

Seien A , B und C beliebige Mengen. Die folgenden Aussagen sind immer wahr:

- (i) $A \subseteq A$
- (ii) $\emptyset \subseteq A$
- (iii) $A \subseteq B \wedge B \subseteq C \Rightarrow A \subseteq C$
- (iv) $A \subseteq B \wedge B \subseteq A \Leftrightarrow A = B$

Manche Menschen haben Schwierigkeiten mit der Aussage, dass die leere Menge Teilmenge jeder Menge A ist. Man kann sich das z. B. folgendermaßen klarmachen: Wäre $\emptyset$ *nicht* Teilmenge von A , dann gäbe es ein Element von $\emptyset$, das nicht Element von A ist. Aber es gibt überhaupt kein Element von $\emptyset$, also ist das nicht möglich.

Punkt (iv) von Lemma 2.1 besagt, dass zwei Mengen *genau dann* gleich sind, wenn sie dieselben Elemente haben. Dass identische Mengen dieselben Element haben, ist selbstverständlich. Dass aber umgekehrt aus der Tatsache, dass zwei Mengen A und B dieselben Elemente haben, bereits folgt, dass die beiden Mengen identisch

sind, ergibt sich daraus, dass eine Menge außer „Ist x ein Element von Dir?“ keine weiteren „Fragen“ beantworten kann.¹⁰ Andere Unterscheidungsmerkmale gibt es also nicht.

Ist A eine Teilmenge von B und gilt $A \neq B$, so nennt man A eine **echte Teilmenge** von B und schreibt $A \subset B$.

Definition

Beispielsweise gilt $\{3, 6\} \subset \{3, 6, 8\}$ und $\{3, 6\} \subseteq \{3, 6\}$, aber *nicht* $\{3, 6\} \subset \{3, 6\}$. Aus $A \subset B$ folgt $A \subseteq B$, aber aus $A \subseteq B$ muss nicht $A \subset B$ folgen.

Aufgabe 2.2. Welche der drei Aussagen $0 = \emptyset$, $\{0\} = \emptyset$ und $0 \in \emptyset$ sind wahr?

Für Mengen A und B definiert man die **Vereinigung** (engl. *union*) $A \cup B$ als die Menge, die aus den Objekten besteht, die in mindestens einer der beiden Mengen A oder B enthalten sind. Der **Durchschnitt** (engl. *intersection*) $A \cap B$ ist die Menge der Objekte, die in beiden Mengen enthalten sind. Die (**mentheoretische**) **Differenz** $A \setminus B$ besteht aus den Elementen, die in A , aber *nicht* in B enthalten sind.

Definition

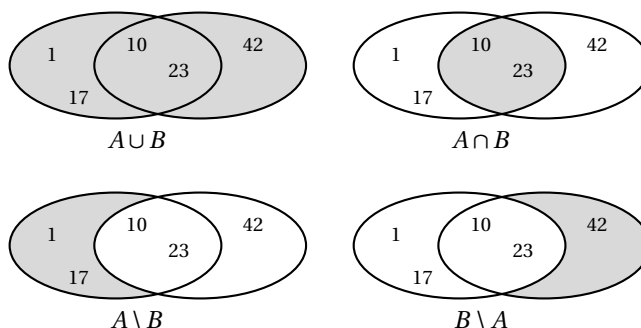
Ein Beispiel mit $A = \{1, 10, 17, 23\}$ und $B = \{10, 23, 42\}$:

$$A \cup B = \{1, 10, 17, 23\} \cup \{10, 23, 42\} = \{1, 10, 17, 23, 42\}$$

$$A \cap B = \{1, 10, 17, 23\} \cap \{10, 23, 42\} = \{10, 23\}$$

$$A \setminus B = \{1, 10, 17, 23\} \setminus \{10, 23, 42\} = \{1, 17\}$$

Natürlich taucht in der ersten Zeile rechts vom Gleichheitszeichen das Element 10 nicht doppelt auf. Wir hatten ja bereits geklärt, dass es keine „mehrfachen Elemente“ gibt. Beachten Sie außerdem, dass in der letzten Zeile das Element 42 keine Rolle spielt. Man kann aus der linken Menge nichts entfernen, was nicht schon drin war. Als Mengendiagramm sieht es folgendermaßen aus:



¹⁰Siehe Seite 19.

Dabei entspricht die grau hinterlegte Fläche jeweils der Menge, die unter der Grafik steht. Beachten Sie, dass im Allgemeinen¹¹ *nicht* $A \setminus B = B \setminus A$ gilt, wie man bereits am allerersten Beispiel sehen kann.

Aufgabe 2.3. Gegeben seien die Mengen $A = \{1, 3, 5, 7\}$, $B = \{2, 3, 6, 7\}$ und $C = \{4, 5, 6, 7\}$. Geben Sie $A \cap B$, $(A \cap B) \cap C$, $A \cup B$, $(A \cup B) \cup C$, $(A \cap B) \cup C$, $(A \cup B) \cap C$, $A \setminus B$, $(A \setminus B) \setminus C$ und $A \setminus (B \setminus C)$ an.

Lemma 2.2

Für alle Mengen A , B und C sind die folgenden Aussagen wahr.

- (i) $A \cap A = A$
- (ii) $A \cup A = A$
- (iii) $A \cap B = B \cap A$
- (iv) $A \cup B = B \cup A$
- (v) $A \cap (B \cap C) = (A \cap B) \cap C$
- (vi) $A \cup (B \cup C) = (A \cup B) \cup C$
- (vii) $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$
- (viii) $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$
- (ix) $A \cup B = A \Leftrightarrow B \subseteq A$
- (x) $A \cap B = A \Leftrightarrow A \subseteq B$
- (xi) $A \cap \emptyset = \emptyset$
- (xii) $A \cup \emptyset = A$
- (xiii) $A \cap B \subseteq A$
- (xiv) $A \subseteq A \cup B$
- (xv) $A \subseteq C \wedge B \subseteq C \Rightarrow A \cup B \subseteq C$
- (xvi) $A \subseteq B \wedge A \subseteq C \Rightarrow A \subseteq B \cap C$
- (xvii) $A \setminus \emptyset = A$
- (xviii) $A \setminus A = \emptyset$
- (xix) $A \setminus (B \cap C) = (A \setminus B) \cup (A \setminus C)$
- (xx) $A \setminus (B \cup C) = (A \setminus B) \cap (A \setminus C)$
- (xxi) $B \subseteq C \Rightarrow A \setminus C \subseteq A \setminus B$
- (xxii) $A \setminus B = \emptyset \Leftrightarrow A \subseteq B$
- (xxiii) $B \cup C \subseteq A \Rightarrow B \setminus C = B \cap (A \setminus C)$

Aufgabe 2.4. Die meisten Aussagen aus Lemma 2.2 sind offensichtlich wahr. Es ist aber trotzdem für Ihr Verständnis sehr wichtig, dass Sie wie bei Lemma 2.1 alle durchgehen und jeweils überlegen, was ausgesagt wird und warum es stimmt. Besonders bei den Aussagen, bei denen mehrere Mengen beteiligt sind, kann es hilfreich sein, mit Mengendiagrammen zu arbeiten. In Aufgabe 2.7 sehen Sie eines für drei Mengen.

¹¹Die Formulierung „im Allgemeinen nicht“ ist typischer Mathematikerschnack und bedeutet, dass es mindestens eine Ausnahme gibt. Man könnte stattdessen auch „nicht immer“ sagen.

Aufgabe 2.5. Die *symmetrische Differenz* zweier Mengen A und B wird durch

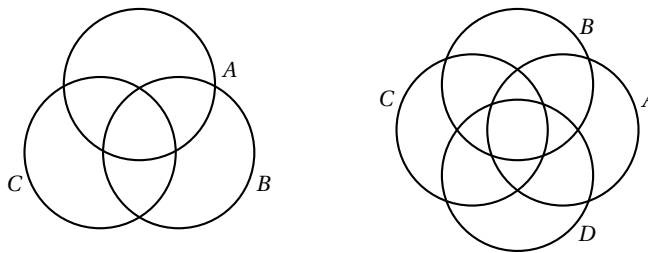
$$A \triangle B = (A \cup B) \setminus (A \cap B) \quad (2.1)$$

symmetrische
Differenz

definiert. Geben Sie $A \triangle B$ und $B \triangle A$ für $A = \{1, 5, 8, 9\}$ und $B = \{8, 1, 12\}$ an.

Aufgabe 2.6. Geben Sie für die symmetrische Differenz eine Alternative zur Definition (2.1) an, in der der Durchschnitt nicht verwendet wird.¹²

Aufgabe 2.7. In der folgenden Skizze sehen Sie links ein Mengendiagramm für drei Mengen.



Der rechte Teil ist jedoch *kein* Mengendiagramm für vier Mengen. Zumindest kein vollständiges in dem Sinne, dass sich jede mögliche Konstellation von vier Mengen so darstellen lässt. Warum nicht?

Aufgabe 2.8. Wie könnte ein vollständiges Mengendiagramm für vier Mengen aussehen? (Die Ränder müssen nicht notwendig Kreise oder Ellipsen sein.)

Bemerkungen

Für das Verständnis ganz wichtig ist die Sichtweise, dass es jedes mathematische Objekt nur einmal gibt, man aber die meisten Objekte auf viele verschiedene Arten hinschreiben kann. Beispielsweise ist $2 + 5$ nur eine andere Art, die Zahl 7 hinzuschreiben. Wenn man $2 + 5 = 7$ schreibt, meint man also, dass links und rechts vom Gleichheitszeichen *dasselbe* steht. Daher ist z. B. $\{7, -1\}$ dieselbe Menge wie $\{2 + 5, 3 - 4\}$. Aus diesem Grund darf in Mengendiagrammen jedes Element auch nur einmal auftreten.

In diesem Zusammenhang sei noch einmal wiederholt, dass x und $\{x\}$ verschiedene Objekte sind. Wenn die Antwort auf eine Frage beispielsweise $\{42\}$ wäre und Sie nur 42 schreiben, so ist Ihre Antwort falsch (weil Sie etwas nicht richtig verstanden haben).

In manchen Büchern wird statt des Kommas das Semikolon als Trennzeichen in Mengen (und eventuell auch in **Tupeln** oder **Intervallen**) verwendet, um Verwechslungen mit dem **Dezimalkomma** zu vermeiden. In diesem Skript wird das nicht gemacht. Der Unterschied ist ansonsten nicht relevant, solange man konsistent immer dieselbe Schreibweise verwendet.

¹²Verweise in Klammern wie (2.1) beziehen sich auf abgesetzte und nummerierte Formeln wie in diesem Fall auf die in Aufgabe 2.5.

So etwas wie $a, b \in M$ ist eine häufig verwendete Abkürzung für $a \in M \wedge b \in M$. Wir werden das später auch machen.

Leider wird in vielen Büchern das Zeichen $\subset$ statt $\subseteq$ benutzt. Ich halte das für einen didaktischen Fehler, aber Sie müssen damit rechnen, dieser Schreibweise öfter zu begegnen.

Die Negation $x \notin M$, die für $\neg(x \in M)$ steht, wird statt „ x ist nicht Element von M “ meistens als „ x ist kein Element von M “ gesprochen oder geschrieben. Das ist eine weit verbreitete Art, Negationen zu formulieren. Man verwendet beispielsweise auch statt der umständlichen Formulierung „ n ist nicht eine Primzahl“ den kürzeren Satz „ n ist keine Primzahl“. Sie sollten sich daran gewöhnen, solche Sätze als logische Negationen zu interpretieren.

Falls Sie eine Eselsbrücke brauchen, weil bei $\cup$ und $\cap$ Verwechslungsgefahr besteht, dann stellen Sie sich vielleicht $\cup$ als einen Topf vor, in den alles hineingeworfen wird.

Machen Sie sich klar, dass die Formulierung „eine leere Menge“ keinen Sinn ergibt. Das wäre so, als würden Sie „ein Eiffelturm“ sagen. Es muss natürlich „*die* leere Menge“ heißen. Es gibt leider immer wieder Studierende, die auch nach zwei Semestern noch „eine leere Menge“ sagen ...

Es gibt keine spezielle Schreibweise für das Hinzufügen bzw. Entfernen eines einzelnen Elementes. Mit den Mengen $A = \{6, 8, 9\}$ und $B = \{6, 7, 8, 9\}$ gilt beispielsweise $A = B \setminus \{7\}$ und $B = A \cup \{7\}$. Es ist *falsch*, so etwas wie $A = B \setminus 7$ oder $B = A \cup 7$ zu schreiben. Symbole wie $\cup$, $\cap$ und $\setminus$ ergeben nur zwischen Mengen Sinn.

Was wir in diesem Skript lernen, ist eine vereinfachte Version der Mengenlehre, die man manchmal auch „naive Mengenlehre“ nennt. Für unsere Zwecke wird das reichen, aber dieser Ansatz führt zu unauflösbaren Widersprüchen, die schon in der obigen Definition des Begriffs *Menge* stecken, der zu vage ist. Die korrekte Herangehensweise ist die sogenannte *axiomatische Mengenlehre*. Wenn Sie mehr darüber wissen wollen, dann schauen Sie sich mein Video [Das Zermelo-Fraenkel-Axiomensystem der Mengenlehre](#) an.

★ Mengen im Computer

PYTHON hat einen eigenen Datentyp `set` für Mengen und man kann „kleine“ Mengen so eingeben, wie man sie auch von Hand schreiben würde:¹³

```
A = {1, 10, 17, 23}
B = {10, 23, 42}
```

Nach den obigen Zuweisungen können Sie z. B. das Folgende ausprobieren:

¹³Lediglich für die leere Menge muss man `set()` schreiben, weil `{}` in PYTHON eine andere Bedeutung hat.

```

10 in A      # Element
{10} <= A    # Teilmenge
A | B        # Vereinigung
A & B        # Durchschnitt
A - B        # mengentheoretische Differenz

```

Vielleicht gehen Sie Aufgabe 2.4 mal mit PYTHON durch?

Man kann auch Elemente zu Mengen hinzufügen bzw. vorhandene entfernen:

```

A.add(100)
B.remove(42)

```

Damit macht man allerdings etwas, was in der Mathematik tabu ist. Wenn dort eine Bezeichnung wie B vergeben wird, dann wird sich diese Zuordnung nicht wieder ändern. Es wird in einem Fachtext *nie* so sein, dass beispielsweise B erst für $\{10, 23, 42\}$ steht und etwas später dann auf einmal für $\{10, 23\}$.¹⁴ Mengen in PYTHON sind *mutable* (veränderlich), mathematische Objekte nicht.

Dadurch werden Mengen im Computer jedoch zur Datenstruktur der Wahl, wenn es darum geht, in einem Programm zu speichern, ob bestimmte Objekte schon einmal untersucht oder bearbeitet wurden. Sie eignen sich dafür besser als beispielsweise *Arrays*.

3. Prädikatenlogik

Mittels der Aussagenlogik kann man bestimmte wichtige Sachverhalte nicht ausdrücken. Dafür muss das Vokabular erweitert werden zur Prädikatenlogik, die auf Vorarbeiten des Deutschen *Gottlob Frege* basiert und Anfang des 20. Jahrhunderts unter anderem durch *David Hilbert*, *Bertrand Russell* und *Alfred Tarski* in die heutige Form gebracht wurde.

Die folgenden Definitionen wirken anfangs wahrscheinlich etwas aufwendig und formal. Das lässt sich nicht vermeiden, wenn man sich präzise ausdrücken will (und eigentlich müsste man noch kleinlicher sein). Im Rahmen eines Informatikstudiums sollte Ihnen das bekannt vorkommen, weil man auch in Programmiersprachen nicht umhin kommt, genau auf die Details zu achten.

Eine **Aussageform** ist ein sprachliches Gebilde, in dem **Individuenvariablen** vorkommen können und das dadurch, dass man für diese Variablen konkrete Werte einsetzt, zu einer Aussage wird.

Definition

„ x ist negativ“ ist z. B. eine Aussageform mit der Individuenvariable x . In dieser Form handelt es sich nicht um eine Aussage, denn es ergibt keinen Sinn, über die

¹⁴Natürlich kann wie im Skript A erst für eine Aussage und dann ein paar Seiten später in einem anderen Abschnitt für eine Menge stehen. Aber so eine „Transformation“ findet nicht innerhalb ein und desselben Gedankengangs statt.

Wahrheit von „ x ist negativ“ zu sprechen. Setzt man für x jedoch die Zahlen -3 oder 42 ein, so ergibt sich in beiden Fällen eine Aussage; im ersten Fall eine wahre, im zweiten eine falsche. Man würde für so eine Aussageform eine Abkürzung wie $A(x)$ verwenden, an der man die Individuenvariable erkennt.¹⁵

Ein paar Anmerkungen dazu:

- Auch so etwas wie $m > n$ ist eine Aussageform, nämlich eine mit *zwei* Individuenvariablen, für die man dann z. B. $A(m, n)$ schreiben würde. Eine Aussage wird daraus erst, indem man für *beide* Variablen Werte einsetzt.
- Aussagen sind auch Aussageformen, und zwar solche mit null Individuenvariablen.
- In einer Aussageform können Variablen auch mehrfach auftreten, beispielsweise in $x \cdot x = x$. Setzt man für so eine Variable einen Wert ein, dann ist das natürlich so gemeint, dass man überall *dasselbe* einsetzt.
- Umgekehrt ist es jedoch nicht verboten, für *verschiedene* Variablen denselben Wert einzusetzen.¹⁶ Zum Beispiel kann man aus $m > n$ durch Einsetzen die wahre Aussage $7 > 3$ machen, aber auch die falsche $4 > 4$.
- In der Regel kann man für die Individuenvariablen nur bestimmte Objekte einsetzen, etwa nur Zahlen oder nur Matrizen. Es wäre z. B. sinnlos, in „ x ist negativ“ für x einen Kreis einzusetzen. Meistens geht das jedoch aus dem Kontext hervor. Die *Quantoren*, die wir gleich einführen, liefern außerdem klare Regeln dafür, was eingesetzt werden darf.

Offensichtlich sind Verknüpfungen von Aussageformen durch Junktoren wieder Aussageformen. Beispielsweise ist $3 < n \wedge n < 8$ eine Aussageform.¹⁷

Aufgabe 3.1. Drücken Sie die Aussageform $x \leq y$ mithilfe der beiden Aussageformen $x < y$ und $x = y$ aus, indem Sie diese durch einen Junktor verbinden.

Aufgabe 3.2. Drücken Sie nun die Aussageform $x < y$ mithilfe der beiden Aussageformen $x \leq y$ und $x = y$ aus.

Aufgabe 3.3. Geben Sie alle Zahlen an, die man in der Aussageform $x \cdot x = x$ für x einsetzen kann, damit daraus eine wahre Aussage wird.

Aufgabe 3.4. Geben Sie eine Aussageform $A(x)$ an, in der man für x Zahlen einsetzen kann und die für *jede* Zahl, die man einsetzt, zu einer wahren Aussage wird. Geben Sie außerdem eine Aussageform $B(x)$ an, die durch *keine* Zahl, die man einsetzen könnte, zu einer wahren Aussage wird.

Aufgabe 3.5. Im 1582 eingeführten *gregorianischen Kalender* ist ein Jahr dann ein *Schaltjahr*, wenn die Jahreszahl durch 4, aber *nicht* durch 100 teilbar ist. Eine Ausnahme von dieser Regel bilden die durch 400 teilbaren Jahreszahlen: die gehören zu

¹⁵Wir werden in diesem Abschnitt für Individuenvariablen kleine Buchstaben wie x und y verwenden und für Aussagen und Aussageformen große wie A und B .

¹⁶Das wurde bereits auf Seite 16 thematisiert.

¹⁷Auch in solchen Fällen wird gerne abgekürzt. Schreibt man beispielsweise $3 < n < 8$, so meint man damit eigentlich die obige Konjunktion.

Schaltjahren. (Beispiele: 1940, 1984 und 2000 sind Schaltjahre; 1900 und 1993 sind keine Schaltjahre.) Wir führen die folgenden Aussageformen ein, wobei die Individuenvariable n jeweils für eine Jahreszahl ab 1582 stehen kann.

$A_4(n)$: n ist durch 4 teilbar.

$A_{100}(n)$: n ist durch 100 teilbar.

$A_{400}(n)$: n ist durch 400 teilbar.

Basteln Sie daraus mithilfe von Junktoren eine Aussageform mit der Individuenvariable n , die genau dann wahr ist, wenn n ein Schaltjahr ist.

Ist $A(\dots)$ eine Aussageform, x eine Individuenvariable und M eine Menge, so sind $\forall x \in M A(\dots)$ und $\exists x \in M A(\dots)$ Aussageformen, in denen x als **gebunden** bezeichnet wird. Variablen, die nicht gebunden sind, werden **frei** genannt. Die beiden neuen Zeichen heißen **Quantoren**; $\forall$ ist der **Allquantor** und $\exists$ der **Existenzquantor**.

Definition

$\forall x \in M A(\dots)$ wird „für alle x aus M gilt $A(\dots)$ “ gelesen, und $\exists x \in M A(\dots)$ als „es gibt ein x aus M , so dass $A(\dots)$ gilt“.

Beispielsweise sind $x < y$, $\forall x \in M (x < y)$ und $\exists y \in M \forall x \in M (x < y)$ drei Aussageformen, von denen die erste zwei freie Variablen hat, die zweite eine und die dritte gar keine. Wenn wir in Zukunft $A(\dots)$ schreiben, so soll damit gemeint sein, dass die zwischen den Klammern stehenden Variablen alle frei sind und es keine weiteren freien Variablen gibt.

Ist $A(x)$ eine Aussageform mit der freien Variable x und m ein konkreter Wert, den man für x einsetzen kann, so schreiben wir $A[m]$ für die Aussage, die sich durch das Einsetzen ergibt. Ist also beispielsweise $A(x)$ die Aussageform $x = 7$, in die man für x Zahlen einsetzen kann, so steht $A[3]$ für die (falsche) Aussage $3 = 7$. Analog gehen wir vor, wenn in einer Aussageform mehrere freie Variablen vorkommen: Ist z. B. $B(x, y)$ die Aussageform $x < y$, dann steht $B[4, 9]$ für die (wahre) Aussage $4 < 9$.

$A[m]$

- Nach den genannten Regeln ist nicht klar, ob die Variable y in der Aussageform $(y = z) \wedge \forall y \in M (y < x)$ frei oder gebunden ist. Solche Fälle werden wir deshalb vermeiden.¹⁸
- Ausdrücke wie $\forall x \in M (y = z)$ oder $\exists x \in M \forall x \in M (x < y)$, in denen ein Quantor mit einer Variablen vor einer Aussageform steht, in der diese Variable gar nicht frei vorkommt, sind nach der Definition nicht verboten. Auch das werden wir jedoch nicht machen. (Sie werden gleich sehen, dass es auch sinnlos wäre.)
- Wenn Sie aufmerksam gelesen haben, dann ist Ihnen aufgefallen, dass wir bei der Einführung der Schreibweise $A[m]$ eigentlich über zwei grundverschiedene Sorten von Variablen sprechen. x ist als Individuenvariable ein

¹⁸Falls es Sie interessiert, wie man mit so einer Situation „professionell“ umgeht: In der mathematischen Logik sind eigentlich nicht die Variablen frei oder gebunden, sondern ihre *Vorkommen* innerhalb eines Ausdrucks. y kommt in der Aussageform dreimal vor. Das y links ist frei, die beiden rechts daneben nicht.

Bestandteil der Formel, während m eine Variable in der Sprache ist, in der wir über die Formel sprechen. (Das ist die sogenannte *Metasprache*.) In ausführlicheren Darstellungen der mathematischen Logik unterscheidet man das, indem man beispielsweise $\overline{m}$ statt m schreibt. Da wir so tief nicht einsteigen wollen, verzichte ich darauf, und hoffe, dass jeweils klar ist, was gemeint ist.

Die Intention der obigen Definition ist, dass jeder Quantor eine Variable bindet, bis keine freien Variablen mehr übrig bleiben. Das Ergebnis ist dann ein Aussage. Wie die zu interpretieren ist, wird durch die folgende Definition festgelegt.

Definition

$\forall x \in M A(x)$ ist eine Aussage, die genau dann wahr ist, wenn $A[m]$ für *jedes* Element m der Menge M wahr ist. Und $\exists x \in M A(x)$ ist eine Aussage, die genau dann wahr ist, wenn $A[m]$ für *mindestens* ein Element m der Menge M wahr ist.

Ist z. B. $M = \{1, 2, 3, 4, 5\}$, dann ist $\exists n \in M (n > 3)$ wahr, weil man für n die Elemente 4 oder 5 von M einsetzen kann: Schreiben wir $A(n)$ für $n > 3$, dann sind $A[4]$ und $A[5]$ wahre Aussagen. $\forall n \in M (n > 3)$ ist jedoch nicht wahr, weil $A(n)$ unter anderem dann nicht wahr ist, wenn man das Element 1 der Menge M für n einsetzt: $A[1]$ ist nicht wahr.

Aufgabe 3.6. Sei $M = \{1, 2, 3, 4, 5\}$, $A(n)$ die Aussageform $n < 3$ und $B(n)$ die Aussageform $n > 3$. Welche von den drei Aussagen $\exists n \in M A(n)$, $\exists n \in M B(n)$ und $\exists n \in M (A(n) \wedge B(n))$ sind wahr?

Aufgabe 3.7. Sei $M = \{1, 2, 3, 4, 5\}$, $A(n)$ die Aussageform $n < 4$ und $B(n)$ die Aussageform $n > 2$. Welche von den drei Aussagen $\forall n \in M A(n)$, $\forall n \in M B(n)$ und $\forall n \in M (A(n) \vee B(n))$ sind wahr?

Kehren wir noch einmal zum Beispiel vor Aufgabe 3.6 zurück. Es ist hoffentlich inzwischen klar geworden, dass $\exists n \in M (n > 3)$ dasselbe bedeutet wie

$$(1 > 3) \vee (2 > 3) \vee (3 > 3) \vee (4 > 3) \vee (5 > 3)$$

und $\forall n \in M (n > 3)$ dasselbe wie

$$(1 > 3) \wedge (2 > 3) \wedge (3 > 3) \wedge (4 > 3) \wedge (5 > 3).$$

Den Existenzquantor kann man sich also allgemein wie eine „große“ Disjunktion vorstellen, den Allquantor wie eine „große“ Konjunktion.¹⁹ Damit sollte auch die folgende Variante der **de-morganschen Gesetze** offensichtlich sein:

Lemma 3.1

De-morgansche Gesetze

$\forall x \in M A(x)$ ist genau dann falsch, wenn $\exists x \in M (\neg A(x))$ wahr ist.

$\exists x \in M A(x)$ ist genau dann falsch, wenn $\forall x \in M (\neg A(x))$ wahr ist.

¹⁹Wenn später unendliche Mengen ins Spiel kommen, kann man das rein formal nicht mehr machen, aber als Analogie ist diese Vorstellung auch dann noch hilfreich.

Der Vollständigkeit halber fügen wir noch eine Regel zur Klammerersparnis (siehe Seite 13) hinzu: Quantoren sollen höhere Priorität als Junktoren haben. Das bedeutet, dass $\forall x \in M A(\dots) \wedge B(\dots)$ beispielsweise eine Abkürzung für die Aussageform $(\forall x \in M A(\dots)) \wedge B(\dots)$ ist. Wir werden jedoch meistens „überflüssige“ Klammern setzen, damit es keine unnötige Verwirrung gibt.

Wir kommen nun zu mehrfachen Quantoren. Sei $A(x, y)$ die Aussageform $x > y$ und M die Menge $\{2, 4, 6, 8, 10\}$. Was besagt die Aussage $\exists x \in M \exists y \in M A(x, y)$? Nach Definition ist sie dann und nur dann wahr, wenn man für x ein Element von M einsetzen kann, das $\exists y \in M A(x, y)$ wahr macht. Probieren wir es aus: Setzen wir dummerweise die Zahl 2 ein, so ergibt sich $\exists y \in M 2 > y$. Hier müssen wir erneut die Definition anwenden, die besagt, dass die Aussage wahr ist, wenn wir für y ein Element aus M einsetzen können, das $2 > y$ wahr macht. Aber das klappt nicht. $2 > 2$ ist nicht wahr, $2 > 4$ ist nicht wahr und so weiter. Aber wir hätten für x ja auch ein anderes Element von M wählen können. Nehmen wir etwa die Zahl 4, so geht es nun um die Aussage $\exists y \in M 4 > y$. Und in diesem Fall können wir $4 > y$ wahr machen, indem wir für y die Zahl 2 einsetzen. Also ist $\exists x \in M \exists y \in M A(x, y)$ wahr. Wir mussten also (mindestens) ein m aus M und (mindestens) ein n aus M finden, so dass $A[m, n]$ wahr wird.

Analog kann man sich überlegen, dass $\forall x \in M \forall y \in M A(x, y)$ *nicht* wahr ist. Es müsste dann nämlich $\forall y \in M A(x, y)$ für jedes Element von M wahr sein, dass man für x einsetzen kann. Das wird aber nie klappen. Beispielsweise ist $\forall y \in M A(6, y)$ nicht wahr, denn das würde ja bedeuten, dass $6 > y$ für jedes Element y von M wahr ist. Aber weder $6 > 6$ noch $6 > 8$ noch $6 > 10$ sind wahr. Damit die Aussage $\forall x \in M \forall y \in M A(x, y)$ wahr ist, muss $A[m, n]$ also *immer* wahr sein, wenn man für m und n irgendwelche Elemente von M einsetzt.

Aufgabe 3.8. Sei $M = \{1, 2, 3\}$ und $N = \{3, 4, 5\}$. Welche der folgenden Aussagen sind wahr und welche falsch?

- (i) $\forall x \in M \forall y \in N x \leq y$
- (ii) $\forall x \in M \forall y \in N x < y$
- (iii) $\exists x \in M \exists y \in N x < y$
- (iv) $\exists x \in N \exists y \in M x \leq y$
- (v) $\exists x \in N \exists y \in M x < y$

Durch die Beispiele ist hoffentlich klar geworden, dass bei gleichen Quantoren die Reihenfolge keine Rolle spielt:

$\exists x \in M \exists y \in N A(x, y)$ ist genau dann wahr, wenn $\exists y \in N \exists x \in M A(x, y)$ wahr ist. $\forall x \in M \forall y \in N A(x, y)$ ist genau dann wahr, wenn $\forall y \in N \forall x \in M A(x, y)$ wahr ist.

Lemma 3.2

Für $\forall x \in M \forall y \in M A(x, y)$ (wenn also auch die beiden Mengen gleich sind), schreibt man daher häufig abkürzend $\forall x, y \in M A(x, y)$. Analog ist die Abkürzung $\exists x, y \in M A(x, y)$ gebräuchlich.

$\forall x, y \dots$
 $\exists x, y \dots$

Bei unterschiedlichen Quantoren dürfen wir jedoch im Allgemeinen nicht einfach vertauschen. Wir führen für ein Beispiel eine Schreibweise ein, die wir in Zukunft noch öfter verwenden werden.

Definition

$a \mid b$ steht dafür, dass a ein Teiler von b ist. Die Negation schreibt man $a \nmid b$.

Zur Auffrischung ein paar Beispiele: 3 ist wegen $21 = 3 \cdot 7$ ein Teiler von 21. Also gilt $3 \mid 21$. Ebenso gilt $1 \mid 3$ und $3 \mid 3$.

Sei nun $M = \{6, 8, 9\}$ und $N = \{2, 3\}$. Die Aussage $\forall x \in M \exists y \in N y \mid x$ ist wahr. Um das zu überprüfen, müssen wir uns nach Definition überzeugen, dass die Aussage $\exists y \in N y \mid x$ für jedes Element von M wahr ist, wenn wir dieses für x einsetzen. Das ist der Fall, denn

- (i) $\exists y \in N y \mid 6$ wird wahr, indem wir 2 oder 3 für y wählen,
- (ii) $\exists y \in N y \mid 8$ wird dadurch wahr, dass wir y durch 2 ersetzen, und
- (iii) $\exists y \in N y \mid 9$ wird durch das Einsetzen von 3 für y wahr.

$\exists y \in N \forall x \in M y \mid x$ ist jedoch *nicht* wahr. Nach Definition würde das bedeuten, dass wir eine Zahl aus N finden, die $\forall x \in M y \mid x$ wahr macht, wenn wir sie für y einsetzen. Aber

- (i) $\forall x \in M 2 \mid x$ ist falsch, weil $2 \nmid 9$ falsch ist, und
- (ii) $\forall x \in M 3 \mid x$ ist falsch, weil $3 \nmid 8$ nicht wahr ist.

Machen Sie sich den Unterschied zwischen diesen beiden Aussagen klar und vergegenwärtigen Sie sich dabei, dass die Reihenfolge, in der die Quantoren aufgeschrieben werden, beim Verständnis hilft: Die erste Aussage besagt, dass wir für jedes x ein y finden müssen, das $y \mid x$ wahr macht. Die zweite Aussage besagt, dass wir ein y finden müssen, das $y \mid x$ für jedes x wahr macht. Im ersten Fall wird zuerst ein Element x von M gewählt. Wir können dann für dieses spezielle x nach einem „passenden“ y suchen. Im zweiten Fall müssen wir ein y finden, das gleichzeitig für *alle* x „passt“. Dass wir die erste Aufgabe erledigen können, heißt noch lange nicht, dass die zweite lösbar ist.

Als Eselsbrücke hilft vielleicht das folgende Beispiel aus dem „wirklichen Leben“. Stellen Sie sich einen Laden vor, der im Schaufenster mit dem Slogan „Für jedes Smartphone die passende Hülle“ wirbt. Ob das wirklich stimmt, müssen Sie wie alle Werbeversprechen nicht glauben, aber es ist zumindest nicht unmöglich. Bezeichnen wir mit S die Menge aller Smartphone-Modelle, mit H die Menge der Hüllen, die das Geschäft im Angebot hat, und mit $P(h, s)$ die Aussage, dass die Hülle h zum Smartphone s passt, dann besagt der Slogan: $\forall s \in S \exists h \in H P(h, s)$. Vertauschen wir die Reihenfolge der Quantoren, so ergibt sich $\exists h \in H \forall s \in S P(h, s)$. Das stimmt nun aber mit Sicherheit nicht, denn es würde bedeuten, dass der Laden (mindestens) eine Hülle im Angebot hat, die für *jedes* Smartphone passt!

Aufgabe 3.9. Sei $M = \{1, 2, 3\}$ und $N = \{2, 3\}$. Welche der folgenden Aussagen sind wahr und welche falsch?

- (i) $\forall x \in M \exists y \in M x = y$
- (ii) $\exists x \in M \forall y \in M x = y$
- (iii) $\forall x \in N \exists y \in M x = y$
- (iv) $\exists y \in M \forall x \in N x = y$
- (v) $\forall x \in M \exists y \in N x = y$
- (vi) $\forall x \in M \exists y \in M x \neq y$
- (vii) $\exists x \in M \forall y \in M x \neq y$
- (viii) $\forall x \in N \exists y \in M x \neq y$
- (ix) $\exists y \in M \forall x \in N x \neq y$
- (x) $\forall x \in M \exists y \in N x \neq y$
- (xi) $\exists y \in N \forall x \in M x \neq y$

Aufgabe 3.10. Wenn M irgendeine Menge ist und $A(x)$ irgendeine Aussageform, so dass $\forall x \in M A(x)$ wahr ist, kann man dann folgern, dass $\exists x \in M A(x)$ wahr ist? Begründen Sie Ihre Antwort.

Aufgabe 3.11. In Aufgabe 1.15 hat man gesehen, dass man mit zwei Junktoren auskommt und die anderen als Abkürzungen interpretieren kann. Könnte man auch auf einen der beiden Quantoren verzichten?

Mit $M = \{8, 9, 55\}$ und $N = \{2, 3, 5, 7\}$ gilt wie im obigen Beispiel $\forall x \in M \exists y \in N y \mid x$. In diesem Fall gilt aber sogar noch ein bisschen mehr: Für jedes Element von M gibt es nur ein einziges Element von N , das man auswählen kann, um die Aussage wahr zu machen. In der Mathematik sagt man dann, dass es *genau ein* Element von M mit der entsprechenden Eigenschaft gibt. Dafür wird die Schreibweise $\exists!$ verwendet,²⁰ in diesem Fall also

genau ein
 $\exists!$

$$\forall x \in M \exists! y \in N y \mid x. \quad (3.1)$$

★**Aufgabe 3.12.** Wie kann man (3.1) durch eine Formel ausdrücken, ohne das neue Symbol $\exists!$ zu verwenden?

Bemerkungen

Der wichtigste Punkt dieses Abschnitts sei hier noch einmal hervorgehoben: Quantoren können dadurch, dass alle freien Variablen gebunden werden, aus Aussageformen Aussagen machen. $n > 1$ ist z. B. eine Aussageform, über deren Wahrheitswert man nichts sagen kann. $\forall n \in \{1, 2\} (n > 1)$ ist jedoch eine Aussage; und zwar eine, die falsch ist.

Eine der gewöhnungsbedürftigen Besonderheiten der mathematischen Fachsprache ist, dass man typischerweise „es gibt ein(e)“ sagt, wenn man im alltäglichen Sprachgebrauch „es gibt *mindestens* ein(e)“ sagen würde. So ist auch der Existenzquantor gemeint. Sie müssen also darauf achten, in solchen Formulierungen das Wort *mindestens* immer mitzudenken, wenn nicht von „genau einem“ Ding die Rede ist.

²⁰In manchen Büchern auch $\exists_1$.

Eine weitere Eigenart der mathematischen Denkweise, die zwar vollkommen logisch ist, zunächst jedoch verwirrend wirken mag, ist die, dass eine Aussage der Form $\forall x \in \emptyset A(x)$ *immer* wahr ist; und zwar völlig unabhängig davon, was $A(x)$ aussagt. Das kann man sich vielleicht am besten mit den de-morganschen Gesetzen (Lemma 3.1) vergegenwärtigen: Wäre $\forall x \in \emptyset A(x)$ falsch, dann wäre $\exists x \in \emptyset \neg A(x)$ wahr. Aber dafür müsste es ein Element m von $\emptyset$ geben, so dass $A[m]$ wahr ist. $\emptyset$ hat jedoch keine Elemente. (Siehe dazu auch Aufgabe 3.10 und den folgenden Abschnitt über Quantoren im Computer.)

Gebundene Variablen sind austauschbar. Damit ist gemeint, dass beispielsweise $\forall x \in M x < 7$ und $\forall y \in M y < 7$ äquivalente Aussagen sind. Allgemein: Tauscht man in einer Aussage, in der eine Variable nur gebunden vorkommt, diese überall (!) gegen eine andere aus, die in der Aussage bisher nicht vorkam, so ändert sich der Wahrheitswert der Aussage dadurch nicht.

★ Quantoren im Computer

Das Pendant zu All- und Existenzquantor sind in PYTHON die Funktionen `all` und `any`. Um für $M = \{-8, 4, 13\}$ zu testen, ob $\forall x \in M x < 10$ wahr ist, gibt man das Folgende ein:

```
M = {-8,4,13}
all(x<10 for x in M)
```

Und $\exists x \in M x^2 < 20$ würde so aussehen:

```
any(x**2<20 for x in M)
```

Denselben Effekt wie `all` könnte man im Prinzip auch so erzielen:

```
answer = True
for x in M:
    if not x<10:
        answer = False
        break
```

Das Ergebnis (True oder False) steht am Ende in der Variablen `answer`. Dadurch wird hoffentlich auch deutlich, warum $\forall x \in \emptyset A(x)$ immer wahr ist. Wie würde so eine Schleife für den Existenzquantor aussehen?

4. „Große“ Mengen

Bisher haben wir nur mit Mengen wie $X = \{23, 42, \pi\}$ gearbeitet, die sehr wenige Elemente haben. Richtig interessant wird es jedoch erst mit „großen“ Mengen. Die kann man allerdings nicht so aufschreiben wie bisher. Manchmal behilft man sich mit **Auslassungspunkten** wie zum Beispiel in

$$Y = \{1, 2, 3, \dots, 100\} \quad \text{oder} \quad Z = \{3, 5, 7, \dots\}, \quad (4.1)$$

aber das ist problematisch. Während es bei Y relativ klar zu sein scheint, dass es um die natürlichen Zahlen von 1 bis 100 geht, kann man bei Z schon ins Grübeln kommen: Ist die nächste Zahl 9, dann sind wohl die ungeraden natürlichen Zahlen gemeint, die größer als 1 sind. Ist sie aber 11, dann sind vielleicht nur die ungeraden Primzahlen gemeint. Oder doch etwas ganz anderes? Eigentlich verlässt man sich darauf, dass die Leser *hoffentlich* verstehen, was man vorhat, denn es sind *immer* mehrere Interpretationen möglich.²¹

Die Lösung dieses Problems besteht darin, eine Eigenschaft zu benennen, die dazu äquivalent ist, dass ein Objekt ein Element der Menge ist. Für die drei obigen Beispiele könnten die Eigenschaften folgendermaßen formuliert werden:

- $x \in X: x = 23 \vee x = 42 \vee x = \pi$
- $x \in Y: x$ ist eine ganze Zahl zwischen 1 und 100
- $x \in Z: x$ ist eine ungerade ganze Zahl, die größer als 1 ist

Das, was jeweils rechts vom Doppelpunkt steht, sind schlicht und einfach Aussageformen, die genau spezifizieren, welche Bedingungen ein Objekt erfüllen muss, damit es ein Element der betreffenden Menge ist. Für derart beschriebene Mengen verwendet man die sogenannte *beschreibende Mengenschreibweise* (engl. *set-builder notation*):

$$M = \{x : A(x)\} \quad (4.2)$$

Dabei ist $A(x)$ eine Aussageform und man liest das als „Menge aller x , für die $A(x)$ gilt“. Anders ausgedrückt ist $x \in M$ für jedes mathematische Objekt x dann und nur dann wahr, wenn $A(x)$ wahr ist. Natürlich kann man statt x auch eine andere Variable verwenden.

Ist z. B. $A = \{1, 2, 3, 4, 5, 6, 7\}$, so ist $\{x : x \in A \wedge 2 \mid x\}$ die Menge $\{2, 4, 6\}$, während $\{k : k \in A \wedge k > 5\}$ die Menge $\{6, 7\}$ ist.

Aufgabe 4.1. Sei $A = \{10, 20, 30, 40, 50\}$. Geben Sie die Mengen $B = \{x : x \in A \wedge x \mid 100\}$ und $C = \{m : m \in A \wedge 4 \mid m\}$ in aufzählender Mengenschreibweise an.

Damit wir zu etwas interessanteren Beispielen als bisher kommen können, führen wir nun nach und nach die wichtigsten Zahlenmengen ein, die wir in den folgenden Kapiteln ohnehin benötigen werden. Die sind Ihnen hoffentlich aus der Schule bekannt, werden jedoch hier noch einmal „offiziell“ eingeführt. Wir beginnen mit einer Menge, die wir mit den uns zur Verfügung stehenden Mitteln nur aufzählend mit Auslassungspunkten schreiben können:

$\mathbb{N}$ steht für die Menge $\{0, 1, 2, 3, 4, \dots\}$. Die Elemente dieser Menge (und nur die) werden **natürlichen Zahlen**, genannt.

Definition

Hierbei sind zwei Dinge zu beachten. Erstens hat $\mathbb{N}$ anders als die bisher betrachteten Mengen²² unendlich viele Elemente. Alles, was wir über Mengen schon gelernt

²¹Mehr dazu in meinem nicht ganz ernst gemeinten Video [Wie man JEDEN Intelligenztest besteht](#).

²²Abgesehen vom Beispiel Z in (4.1).

haben, stimmt weiterhin, aber es erfordert etwas mehr mentale Arbeit, sich unendliche Mengen vorzustellen. Zweitens ist in diesem Skript per definitionem die Null eine natürliche Zahl. Das ist etwas, worüber in der Mathematik keine Einigkeit besteht.²³ Das ist nicht schlimm, man muss sich nur auf eine Sprachregelung einigen. Die Regelung, die für die aktuelle Vorlesung gilt, kennen Sie nun.

Mit der Formulierung „und nur die“ ist gemeint, dass wir ein mathematisches Objekt, das kein Element von $\mathbb{N}$ ist, *nicht* als natürliche Zahl bezeichnen. Dieser Zusatz wird bei den folgenden Mengen als selbstverständlich vorausgesetzt und nicht jedesmal wiederholt.

Da man häufig auch die Menge der positiven natürlichen Zahlen benötigt, führen wir dafür ebenfalls ein Symbol ein: $\mathbb{N}^+ = \mathbb{N} \setminus \{0\}$.

Aufgabe 4.2. Schreiben Sie (unter Verwendung der Menge $\mathbb{N}$) die Menge $\mathbb{N}^+$ in beschreibender Mengenschreibweise auf.

Aufgabe 4.3. Schreiben Sie die beiden Mengen aus (4.1) in beschreibender Mengenschreibweise auf.

Aufgabe 4.4. Schreiben Sie die Menge $\{0, 2, 4, 6, \dots\}$ der geraden natürlichen Zahlen in beschreibender Mengenschreibweise auf.

Definition

Eine **Primzahl** ist eine natürliche Zahl, die größer als 1 ist und die außer 1 und sich selbst keine weiteren positiven Teiler hat. Für die Menge aller Primzahlen wird das Symbol $\mathbb{P}$ verwendet. Eine natürliche Zahl, die größer als 1, aber keine Primzahl ist, wird **zusammengesetzt** (engl. *composite*) genannt.

prim

Für die Eigenschaft einer Zahl, eine Primzahl zu sein, ist auch das Adjektiv *prim* gebräuchlich. Die ersten zehn Primzahlen sind 2, 3, 5, 7, 11, 13, 17, 19, 23 und 29. Man würde also beispielsweise sagen, dass 7 prim ist. Beachten Sie, dass die natürlichen Zahlen 0 und 1 nach Definition weder Primzahlen noch zusammengesetzte Zahlen sind.²⁴

Aufgabe 4.5. Geben Sie für $\mathbb{P}$ und für die Menge der zusammengesetzten Zahlen Schreibweisen ohne Auslassungspunkte an.

Definition

Für $n \in \mathbb{N}$ definieren wir $[n]$ als die Menge $\{k : k \in \mathbb{N}^+ \wedge k \leq n\}$.

Beachten Sie, dass wir hier, was noch öfter geschehen wird, gleich unendlich viele Mengen auf einmal definiert haben – unter anderem $[7]$, $[13]$ und $[2000]$.

Aufgabe 4.6. Verständnisfrage: Welche Mengen sind nach dieser Definition mit $[3]$, $[0]$ und $[\pi]$ gemeint?

²³In meinem Video *Ist eins eine Primzahl?* gehe ich ausführlich darauf ein.

²⁴In dem schon erwähnten Video *Ist eins eine Primzahl?* geht es, wie der Titel bereits andeutet, auch um diese Frage.

Sei $X = \{n : n \in [100] \wedge 2 \mid n\}$, also in aufzählender Schreibweise $\{2, 4, 6, \dots, 100\}$. Diese Menge könnte man auch als

$$\{2 \cdot 1, 2 \cdot 2, 2 \cdot 3, \dots, 2 \cdot 50\}$$

schreiben, wodurch gewissermaßen deutlich wird, wie sie „konstruiert“ wird: Man durchläuft die natürlichen Zahlen von 1 bis 50, also die Elemente von $[50]$, multipliziert jede von denen mit 2 und sammelt die Produkte in der Menge X . Es gibt eine erweiterte Version der beschreibenden Mengenschreibweise, mit der man so etwas darstellen kann. Im konkreten Fall wäre das

$$X = \{2n : n \in [50]\}$$

und in allgemeiner Form

$$M = \{t(x) : A(x)\}, \quad (4.3)$$

wobei $A(x)$ wieder eine Aussageform ist und $t(x)$ ein **Term**, der angibt, wie ein Element von M aussieht, wenn x die Bedingung $A(x)$ erfüllt. Gelesen wird das als „Menge aller $t(x)$, für die $A(x)$ gilt“ und die Vorstellung ist, dass jedes Objekt x , für das $A(x)$ wahr ist, ein Element $t(x)$ zur Menge M beiträgt. $t(x)$ ist quasi ein „Prototyp“ für die Elemente von M . Die beschreibende Mengenschreibweise aus (4.2) ist lediglich ein Spezialfall von (4.3), bei dem $t(x)$ einfach der Term x ist.

Aufgabe 4.7. Geben Sie $A = \{n^2 : n \in [7]\}$ in aufzählender Schreibweise an. Versuchen Sie dann, A im Stil von (4.2) – also ohne Verwendung der eben eingeführten Schreibweise – aufzuschreiben.

Aufgabe 4.8. Geben Sie die Menge aus Aufgabe 4.4 in der neuen Schreibweise an.

Aufgabe 4.9. Geben Sie $\{1, 4, 7, 10, 13, \dots\}$ in beschreibender Mengenschreibweise an.

Aufgabe 4.10. Geben Sie die beiden Mengen

$$A = \{x^2 : x \in \mathbb{N} \wedge x < 6\} \quad \text{und} \quad B = \{x : x \in \mathbb{N} \wedge x^2 < 6\}$$

in aufzählender Mengenschreibweise an.

$\mathbb{Z} = \mathbb{N} \cup \{-n : n \in \mathbb{N}\}$ ist die Menge der **ganzen Zahlen** (engl. *integers*).

Definition

Bei dieser Definition kommt die Null „doppelt“ vor. Aber das macht ja nichts, weil Mengen solche Dopplungen ignorieren. Das international übliche Symbol $\mathbb{Z}$ basiert übrigens auf dem deutschen Wort *Zahl*.

Aufgabe 4.11. Geben Sie $M = \{m : m \in \mathbb{Z} \wedge m^2 < 6\}$ in aufzählender Mengenschreibweise an.

Aufgabe 4.12. Finden Sie für die folgenden Mengen jeweils noch eine andere Darstellung. Auslassungspunkte sind dabei nicht zugelassen.

$$A = \{n - 3 : n \in \mathbb{N}\}$$

$$B = \{z - 13 : z \in \mathbb{Z}\}$$

$$C = \mathbb{N}^+ \setminus [3]$$

$$D = C \cup \{-k : k \in C\}$$

$$E = \{3n + 1 : n \in \mathbb{N}\} \cup \{3n + 2 : n \in \mathbb{N}\}$$

Die Darstellung (4.3) lässt sich noch weiter verallgemeinern, indem man mehrere Variablen zulässt. Ein Beispiel dafür ist

$$X = \{a - b : a \in A \wedge b \in B\} \quad \text{mit}$$

$$A = \{10, 20, 30\} \quad \text{und} \quad B = \{1, 2, 3, 4\}.$$

Das ist so gemeint, dass die Variablen a und b *unabhängig voneinander* Werte annehmen können, solange die Aussageform rechts vom Doppelpunkt wahr ist. Diese Werte werden dann zu dem Ausdruck links vom Doppelpunkt kombiniert, woraus sich ein Element der Menge X ergibt. Z. B. kann a den Wert 10 haben und b den Wert 2; dann ergibt sich $a - b = 8$ und damit ist 8 ein Element von X . Mit $a = 20$ und $b = 2$ bekommt man das Element 18, mit $a = 10$ und $b = 4$ das Element 6 und so weiter. Insgesamt hat man damit in diesem Beispiel.

$$X = \{6, 7, 8, 9, 16, 17, 18, 19, 26, 27, 28, 29\}.$$

Die allgemeine Form ist

$$\{t(x_1, \dots, x_n) : A(x_1, \dots, x_n)\}, \quad (4.4)$$

wobei $t(x_1, \dots, x_n)$ ein Term ist, der von den n Variablen x_1 bis x_n abhängt. Dieser Term ist nach wie vor ein „Prototyp“ für die Elemente der definierten Menge. Im obigen Beispiel handelte es sich um den Term $a - b$ mit den zwei Variablen a und b .

Aufgabe 4.13. Geben Sie die Menge $\{mn : m, n \in [4]\}$ in aufzählender Mengenschreibweise an.

Definition

$\mathbb{Q} = \{\frac{p}{q} : p \in \mathbb{Z} \wedge q \in \mathbb{N}^+\}$ ist die Menge der **rationalen Zahlen**, die auch als **Brüche** (engl. *fractions*) bezeichnet werden.

Diese Definition führt jeden Bruch unendlich oft auf, weil man beispielsweise mit $p = 2$ und $q = 3$ dieselbe Zahl wie mit $p = 10$ und $q = 15$ erhält. Wie bereits bei der Definition von $\mathbb{Z}$ sind diese Dopplungen jedoch unproblematisch. Das hier verwendete Symbol $\mathbb{Q}$ basiert auf dem Wort *Quotient*, weil Brüche durch Division von ganzen Zahlen entstehen.

Aufgabe 4.14. Finden Sie für die folgenden Mengen jeweils noch eine andere Darstellung. Auslassungspunkte sind dabei nicht zugelassen.

$$A = \left\{ \frac{m}{n} : m \in \mathbb{N} \wedge n \in \mathbb{N}^+ \right\}$$

$$B = \left\{ \frac{m}{n} : m \in \mathbb{N} \wedge n \in \mathbb{N}^+ \wedge m < n \right\}$$

$$C = \{ q : q \in \mathbb{Q} \wedge q^2 < 1 \}$$

$$D = \{ -z : z \in \mathbb{Z} \}$$

$$E = \{ 5q : q \in \mathbb{Q} \}$$

Schließlich noch eine letzte Variante der beschreibenden Mengenschreibweise. Ihnen wird aufgefallen sein, dass auf der rechten Seite des Doppelpunktes fast immer so etwas wie $x \in M$ vorkommt. Es ist erlaubt, diesen Teil der Bedingung vor den Doppelpunkt zu ziehen, wenn links vom Doppelpunkt nur eine Variable stand. Beispielsweise darf man statt $\{ n : n \in \mathbb{N} \wedge n < 10 \}$ auch $\{ n \in \mathbb{N} : n < 10 \}$ schreiben. Das ist keine erwähnenswerte Abkürzung, da man kaum etwas spart, aber es hilft beim schnellen Erfassen der Menge, weil man sofort sieht, aus welcher Grundmenge die „Kandidaten“ entnommen werden.

Die allgemeine Form ist

$$\{ x \in M : A(x) \} \tag{4.5}$$

also eine Variante von $\{ x : x \in M \wedge A(x) \}$. Man liest das als „Menge der x aus M , für die die Bedingung $A(x)$ gilt“.

Aufgabe 4.15. Verwenden Sie die Schreibweise (4.5) für die Definitionen von $\mathbb{N}^+$ (wie in Aufgabe 4.2), für $[n]$ und für die Menge X der ganzen Zahlen, die kleiner als 42 sind.

Aufgabe 4.16. Für welche der beiden Mengen aus Aufgabe 4.10 kann man die Schreibweise (4.5) verwenden?

Aufgabe 4.17. Für $n \in \mathbb{N}$ definieren wir $A_n = \{ k \in \mathbb{N} : n \leq k \leq 2n \}$. Geben Sie A_0 , A_3 und A_5 in aufzählender Schreibweise an.

Aufgabe 4.18. Welche der folgenden Aussagen sind wahr?

- (i) $\exists n \in \mathbb{N} (2 \nmid n \wedge 4 \mid n)$
- (ii) $\exists n \in \mathbb{N} (2 \mid n \wedge 4 \nmid n)$
- (iii) $(\forall n \in \mathbb{N} 2 \mid n) \Rightarrow (\forall n \in \mathbb{N} 4 \mid n)$
- (iv) $\forall n \in \mathbb{N} (2 \mid n \Rightarrow 4 \mid n)$
- (v) $\forall n \in \mathbb{N} (4 \mid n \Rightarrow 2 \mid n)$
- (vi) $\forall n \in \mathbb{N} (2 \mid n \Leftrightarrow 4 \mid n)$
- (vii) $\forall n \in \mathbb{N} (4 \nmid n \Rightarrow 2 \nmid n)$

Aufgabe 4.19. Welche der folgenden Aussagen sind wahr?

- (i) $\forall m \in \mathbb{N} \exists n \in \mathbb{N} n > m$
- (ii) $\exists n \in \mathbb{N} \forall m \in \mathbb{N} n > m$
- (iii) $\forall m \in \mathbb{N} \exists n \in \mathbb{N} mn = n$

$$(iv) \exists n \in \mathbb{Z} \forall m \in \mathbb{Z} mn = n$$

Aufgabe 4.20. Mit der beschreibenden Mengenschreibweise hätte man die Vereinigung zweier Mengen auch als $A \cup B = \{x : x \in A \vee x \in B\}$ definieren können. Wie müssten Durchschnitt und Differenz von zwei Mengen definiert werden?

Aufgabe 4.21. Und wie kann man die **symmetrischen Differenz** von A und B mit der beschreibenden Mengenschreibweise darstellen, ohne $\cap$, $\cup$ und $\setminus$ zu verwenden?

Aufgabe 4.22. Schauen Sie sich die Punkte (xix) und (xx) von Lemma 2.2 noch einmal an und machen Sie sich klar, dass es sich dabei um die bereits mehrfach thematisierten **de-morganschen Gesetze** handelt.

Aufgabe 4.23. Seien $A(x)$ und $B(x)$ Aussageformen, so dass $\forall x \in X (A(x) \Rightarrow B(x))$ für eine Menge X gilt. Wenn M und N definiert sind durch

$$M = \{x \in X : A(x)\} \quad \text{und} \quad N = \{x \in X : B(x)\},$$

welcher Zusammenhang besteht dann zwischen diesen beiden Mengen?

Aufgabe 4.24. Es ist wohl offensichtlich, dass $\mathbb{P} \subseteq \mathbb{N} \subseteq \mathbb{Z} \subseteq \mathbb{Q}$ gilt. Begründen Sie, dass man an allen drei Stellen $\subseteq$ durch $\subset$ ersetzen kann, indem Sie jeweils eine Zahl angeben, die in der Obermenge, aber nicht in der Teilmenge enthalten ist. (Und ist Ihnen wirklich klar, warum nach den Definitionen in diesem Abschnitt $\mathbb{Z} \subseteq \mathbb{Q}$ gelten muss? Ansonsten schauen Sie bitte in der Lösung nach.)

Aufgabe 4.25. A sei die Menge der natürlichen Zahlen, deren Quadrat kleiner als 400 und durch 12 teilbar ist. B sei die Menge der natürlichen Zahlen, deren Quadrat kleiner als 400 ist und die durch 12 teilbar sind. Geben Sie die beiden Mengen in beschreibender Mengenschreibweise an. Gilt $A = B$?

Bemerkungen

Variablen werden durch die beschreibende Mengenschreibweise wie durch Quantoren gebunden und sind daher ebenfalls austauschbar. Das bedeutet, dass z. B.

$$\{n : n \in \mathbb{N} \wedge n < 100\} \quad \text{und} \quad \{k : k \in \mathbb{N} \wedge k < 100\}$$

dieselbe Menge beschreiben. Wie in der Prädikatenlogik müssen in so einem Fall natürlich erstens *alle* Vorkommen einer Variablen ersetzt werden und zweitens darf die neue Variable keine sein, die in der Mengendefinition bereits vorkam.

Die Aussageform rechts vom Doppelpunkt in der beschreibenden Mengenschreibweise muss immer ausführlich genug sein, damit ohne Zweifel klar ist, was zur Menge gehört und was nicht. $\{x : x^2 < 6\}$ würde z. B. nicht reichen. Man vergleiche dazu die Mengen B in Aufgabe 4.10 und M in Aufgabe 4.11. In der Regel wird für eine korrekte Beschreibung mindestens eine bereits bekannte Menge wie etwa $\mathbb{N}$ oder $\mathbb{Q}$ benötigt.

Die Division zweier Zahlen x und y , die in der Schule oft die Form $x : y$ hat, wird in der Hochschulmathematik meistens als x/y oder als $\frac{x}{y}$ geschrieben. Es ist kein

Problem, dass das wie ein Bruch aussieht, denn der Bruch p/q ist ja das Ergebnis der Division von p durch q .

Beachten Sie, dass die Schreibweise (4.5) nur als Variante von (4.2) angewendet werden sollte. Sie sollte nicht mit (4.4) vermischt werden, wenn $t(x)$ etwas anderes als nur x ist, weil dann evtl. nicht mehr klar ist, was gemeint sein soll. So etwas wie $\{x - 7 \in \mathbb{N} : x < 10\}$ ist z. B. *falsch*! Gehört -4 zur Menge, weil $x = 3$ kleiner als 10 ist, oder gehört -4 nicht zur Menge, weil $3 - 7$ kein Element von $\mathbb{N}$ ist?

Manchmal sieht man in Klausuren Schreibweisen wie $\mathbb{Z} = \{n, -n : n \in \mathbb{N}\}$. Man kann sich zwar denken, was damit gemeint ist, aber diese Schreibweise ist *falsch*! Links vom Doppelpunkt steht immer nur *ein* Prototyp.

In einigen Büchern wird statt des Doppelpunktes ein senkrechter Strich verwendet: $\{x \mid A(x)\}$. Damit ist dasselbe wie hier gemeint.²⁵ Es ist egal, welche Konvention man verwendet, solange man es konsistent handhabt.

Wir kennen bereits Abkürzungen wie $x, y \in M$ für $x \in M \wedge y \in M$. Häufig wird auch in der Mengenschreibweise ein Komma statt $\wedge$ verwendet, das könnte dann z. B. die Form $\mathbb{Q} = \{p/q : p \in \mathbb{Z}, q \in \mathbb{N}^+\}$ annehmen.

Das Adjektiv *positiv* bedeutet „größer als null“. Die Zahl null ist also *nicht* positiv. Das Adjektiv *negativ* steht für „kleiner als null“. Die Null ist also auch nicht negativ. Will man $x \geq 0$ ausdrücken, so sagt man in der Mathematik, dass x *nichtnegativ* ist. Mehr dazu in Abschnitt 11.

positiv
negativ
nichtnegativ

★ Mengenschreibweise im Computer

PYTHON bietet eine *comprehension* genannte Syntax, die der beschreibenden Mengenschreibweise sehr ähnlich ist und von dieser inspiriert wurde.²⁶ Sie wurde schon im [Abschnitt über Quantoren im Computer](#) verwendet. Hier noch ein paar einfache Beispiele:

```
A = {1, 10, 17, 23}
B = {10, 23, 42}
C = {x for x in A if x**2 > 100}
D = {x**2 for x in B}
E = {x*y for x in A for y in B}
```

Damit hat man $C = \{x \in A : x^2 > 100\}$, $D = \{x^2 : x \in B\}$ und $E = \{xy : x \in A \wedge y \in B\}$ definiert.

5. Arithmetik

Arithmetik ist das Teilgebiet der Mathematik, das sich mit dem Rechnen beschäftigt. In diesem Zusammenhang gehen wir etwas ausführlicher auf die verschiedenen

²⁵In diesem Zusammenhang ist $\mid$ also *nicht* das Zeichen für Teilbarkeit.

²⁶Man findet diese angenehme Syntax nur in wenigen anderen Programmiersprachen wie [JULIA](#) oder [HASKELL](#).

„Sorten“ von Zahlen ein. Das ist unter anderem auch für die Informatik relevant, weil die meisten Programmiersprachen erstens unterschiedliche Zahlentypen anbieten und zweitens nicht alle Zahlen exakt darstellen können. Größtenteils handelt es sich bei den folgenden Seiten um Wiederholung von Schulstoff, der hier lediglich zusammengefasst und präzisiert wird.

Satz 5.1

Summen und Produkte von natürlichen Zahlen sind wieder natürliche Zahlen. Für alle natürlichen Zahlen a , b und c gilt:

- (i) $a + b = b + a$
- (ii) $(a + b) + c = a + (b + c)$
- (iii) $ab = ba$
- (iv) $(ab)c = a(bc)$
- (v) $a(b + c) = ab + ac$

Ferner gilt $n + 0 = n$ und $1 \cdot n = n$ für alle natürlichen Zahlen n und die Zahlen 0 und 1 sind die einzigen Zahlen mit diesen Eigenschaften.

Kommutativität

Assoziativität

Distributivität

neutrales Element

Die Eigenschaften (i) und (iii) nennt man die *Kommutativität* der Addition bzw. Multiplikation, bei (ii) und (iv) handelt es sich um die sogenannte *Assoziativität*. Eigenschaft (v) ist die einzige im obigen Satz, die einen Zusammenhang zwischen Addition und Multiplikation herstellt. Sie wird *Distributivität* genannt.

Der letzte Absatz besagt, dass die Null das *neutrale Element* der Addition ist. Entsprechend ist die Eins das neutrale Element der Multiplikation.

Die Null spielt auch bezüglich der Multiplikation eine wichtige Rolle:

Korollar 5.2

Für alle natürlichen Zahlen gilt $0 \cdot n = 0$. Sind ferner m und n natürliche Zahlen mit $mn = 0$, dann muss mindestens einer der beiden Faktoren die Null sein.

Dieses Korollar zeigt eine von vielen sprachlichen Varianten, den Allquantor auszudrücken. Wenn ein Satz mit so etwas wie „Sind m und n natürliche Zahlen“ anfängt, ist damit eigentlich „Für alle $m, n \in \mathbb{N}$ “ gemeint. Formal kann man es daher so schreiben:

$$\forall m, n \in \mathbb{N} (mn = 0 \Rightarrow m = 0 \vee n = 0).$$

Aufgabe 5.1. Die Aussagen von Satz 5.1 kann man alle auch als Formeln der Prädikatenlogik aufschreiben, z. B. $\forall a, b \in \mathbb{N} a + b \in \mathbb{N}$. Machen Sie das zur Übung.

Satz 5.3

Alle Aussagen von Satz 5.1 und Korollar 5.2 gelten auch, wenn man überall „natürliche Zahlen“ durch „ganze Zahlen“ ersetzt. Zusätzlich gibt es nun zu jeder ganzen Zahl a genau eine ganze Zahl b , für die $a + b = 0$ gilt.

inverses Element

$a - b$

Für die Zahl b aus Satz 5.3 schreibt man $-a$. Sie wird das *inverse Element* von a bezüglich der Addition genannt. Ein Ausdruck wie $a - b$ ist in diesem Sinne eine

Abkürzung für $a + (-b)$. (Daraus folgt sofort, dass $a - b$ eine ganze Zahl ist, wenn a und b ganze Zahlen sind.)

Natürlich „passt“ diese Schreibweise zu der Art und Weise, wie wir negative ganze Zahlen schreiben. -5 ist beispielsweise das inverse Element von 5 . Umgekehrt ist jedoch auch 5 das inverse Element von -5 .

Aufgabe 5.2. Verständnisfrage: Was ist das inverse Element von null?

Aufgabe 5.3. Falls Sie die oben besprochene Unterscheidung von inversem Element und Subtraktion übertrieben kleinlich finden, dann überlegen Sie sich, was ein **Compiler** für eine Sprache wie **JAVA** machen muss, wenn er Ausdrücke wie -3 , $-a$ oder $a-b$ liest. In allen drei Fällen kommt das Minuszeichen vor, aber es hat jeweils eine etwas andere Bedeutung, die entsprechend auch zu einer anderen Aktion führen muss.

Man spricht häufig von den „vier Grundrechenarten“. Für das mathematische Verständnis ist es jedoch hilfreich, wenn man sich klarmacht, dass es eigentlich nur *zwei* grundlegende Rechenoperationen gibt, nämlich Addition und Multiplikation. Subtraktion und Division sind in diesem Sinne lediglich abgeleitete Operationen. Das mag Ihnen wie eine unnötige Komplikation vorkommen, aber es ist im Gegenteil eine Denk- und Rechenhilfe. Beispielsweise ist die Subtraktion weder **kommutativ** noch **assoziativ**: $12 - 8 \neq 8 - 12$ und $(3 - 4) - 5 \neq 3 - (4 - 5)$. Fasst man sie jedoch als Addition auf, so kann man die entsprechenden Rechengesetze anwenden:

$$12 - 8 = 12 + (-8) = -8 + 12$$

$$(3 - 4) - 5 = (3 + (-4)) + (-5) = 3 + ((-4) + (-5))$$

Geht man derart vor, dann kann man sofort ein paar einfache Rechenregeln für das Minuszeichen herleiten:

Für alle ganzen Zahl a , b und c gilt:

- (i) $-(-a) = a$
- (ii) $0 - a = -a$
- (iii) $(-a) \cdot b = a \cdot (-b) = -ab$
- (iv) $(-a) \cdot (-b) = ab$
- (v) $a \cdot (b - c) = ab - ac$
- (vi) $(b - c) \cdot a = ba - ca$
- (vii) $-(a - b) = b - a$
- (viii) $a - (b - c) = a - b + c$
- (ix) $a - (b + c) = a - b - c$

Lemma 5.4

Aufgabe 5.4. Berechnen Sie ohne Taschenrechner:

- (i) $(-3) + (-2)$
- (ii) $(-3) \cdot (-2)$

- (iii) $6 \cdot 13 + 4 \cdot 13$
- (iv) $6 + 4 \cdot 10$
- (v) $(4 + 12) + (8 + 6)$
- (vi) $(-1) \cdot (-2) \cdot (-3)$
- (vii) $(-1) \cdot (-1) \cdot (-1) \cdot (-1) \cdot (-1) \cdot (-1)$
- (viii) $100 - (80 - 20)$
- (ix) $(8 - 10) \cdot (5 - 10)$
- (x) $72 - (12 - 20)$
- (xi) $18 \cdot 11 - 11 \cdot 9$

Aufgabe 5.5. Welche der folgenden Aussagen sind wahr?

- (i) $\forall a, b \in \mathbb{Z} \ 6a - 6b = a - b$
- (ii) $\forall a, b \in \mathbb{Z} \ 8ab - ba = 8$
- (iii) $\forall a, b \in \mathbb{Z} \ 2 \cdot (a \cdot b) = 4ab$
- (iv) $\forall a, b, c \in \mathbb{Z} \ a \cdot (b + c) = b \cdot a + c \cdot a$
- (v) $\forall a \in \mathbb{Z} \ 2a + a \cdot 3 = 9a$
- (vi) $\forall b \in \mathbb{Z} \ b + (3b + 4b) = 9b$
- (vii) $\forall a \in \mathbb{Z} \ 3a - a = 3$
- (viii) $\forall a \in \mathbb{Z} \ -(-a) = a$
- (ix) $\forall a, b, c \in \mathbb{Z} \ a + b \cdot c = (a + b) \cdot (a + c)$
- (x) $\forall a, b \in \mathbb{Z} \ -(a - b) = b - a$

Satz 5.5

Alle bisherigen Aussagen gelten auch, wenn man überall „ganze Zahlen“ durch „rationale Zahlen“ ersetzt. Zusätzlich gibt es nun zu jeder rationalen Zahl a außer 0 genau eine rationale Zahl b , für die $ab = 1$ gilt.

Kehrwert

Für die Zahl b aus Satz 5.5 schreibt man $1/a$ oder a^{-1} . Sie wird das *inverse Element* von a bezüglich der Multiplikation oder auch der *Kehrwert* von a genannt. Beachten Sie die Analogie zum inversen Element bezüglich der Addition. Dort soll sich 0 ergeben, hier 1; in beiden Fällen also das neutrale Element der jeweiligen Verknüpfung.

a/b

Ein Ausdruck wie a/b ist als Abkürzung für $a \cdot (1/b)$ gemeint. (Und damit ist klar, dass a/b eine rationale Zahl ist, wenn a und b rationale Zahlen sind.) Für a/b schreibt man auch $\frac{a}{b}$. Wie bereits erwähnt werden sowohl die „Schulschreibweise“ $a : b$ als auch die auf Taschenrechnern und im anglophonen Raum übliche Variante $a \div b$ in der Hochschulmathematik so gut wie nie verwendet.

Hier gelten dieselben Anmerkungen wie bei der Subtraktion: Division ist keine eigenständige Rechenoperation, sondern abgeleitet aus der Multiplikation.

Weil Korollar 5.2 auch für rationale Zahlen gilt, kann die Zahl null kein inverses Element haben. Daher ergeben auch Ausdrücke wie „ $a/0$ “ keinen Sinn.²⁷

²⁷Division durch null ist also nicht „verboten“, sondern es geht einfach nicht. In meinem Video [Warum kann man nicht durch null teilen?](#) gehe ich darauf noch etwas ausführlicher ein.

Für alle rationalen Zahl a, b, c und d gilt:

- (i) $1/(1/a) = a$
- (ii) $1/(a/b) = b/a$
- (iii) $(a/b) \cdot (c/d) = (ac)/(bd)$
- (iv) $(a/b)/(c/d) = (a/b) \cdot (d/c) = (ad)/(bc)$
- (v) $(ac)/(bc) = a/b$
- (vi) $a/b + c/b = (a+c)/b$
- (vii) $a/b + c/d = (ad)/(bd) + (bc)/(bd) = (ad+bc)/(bd)$

Diese Regeln sind selbstverständlich alle so zu verstehen, dass nirgends durch null dividiert wird.

Lemma 5.6

Aufgabe 5.6. Berechnen Sie die folgenden Ausdrücke ohne Taschenrechner. Ein Ergebnis, das nicht vollständig gekürzt ist, gilt in einer Prüfung als falsch!

- (i) $1/15 + 4/15$
- (ii) $5/8 + 3/2$
- (iii) $5/6 + 2/15$
- (iv) $3/2 - 3/4$
- (v) $1/6 - 1/7$

Aufgabe 5.7. Vereinfachen Sie die Ausdrücke

$$\frac{a}{6} + \frac{a}{9}, \quad \frac{a}{10} + \frac{b}{5} \quad \text{und} \quad \frac{a-b}{2} + \frac{b-a}{3}$$

in dem Sinne, dass es danach nur noch einen Bruchstrich gibt. a und b sind Platzhalter für beliebige rationale Zahlen.

Aufgabe 5.8. Vereinfachen Sie die Terme

$$\frac{cm+2c}{dc}, \quad \frac{ax+bx}{x-cx}, \quad \frac{8z}{3} - \frac{5z}{2} \quad \text{und} \quad \frac{x}{3} + \frac{a+x}{2} - \frac{5x+a}{6}.$$

Die Buchstaben sind Platzhalter für beliebige rationale Zahlen, die so gewählt sind, dass nicht durch null geteilt wird.

Aufgabe 5.9. Begründen Sie mit einem konkreten Gegenbeispiel, dass die „Regel“

$$\frac{a}{b} + \frac{c}{d} = \frac{a+c}{b+d}$$

nicht für alle $a, b, c, d \in \mathbb{Q}$ gilt. (Das ist ein „beliebter“ Fehler.)

Aufgabe 5.10. Welche der Gleichungen

$$\begin{aligned} \frac{a+c}{b+c} &= \frac{a}{b}, & \frac{a+c}{b+c} &= \frac{a+1}{b+1}, & \frac{ax+by}{x+y} &= a+b, & \frac{a}{b} + \frac{c}{d} &= \frac{a+b}{c+d}, \\ \frac{1}{a+b} &= \frac{1}{a} + \frac{1}{b}, & \frac{3a-4c}{3x-4y} &= \frac{a-c}{x-y}, & \frac{a}{b} + 1 &= \frac{a+1}{b}, & c \cdot \frac{a}{b} &= \frac{ca}{cb} \end{aligned}$$

sind für alle rationalen Zahlen a, b, c, d, x und y wahr (sofern nicht durch null dividiert wird)? Überlegen Sie sich bei den falschen Gleichungen jeweils ein konkretes Gegenbeispiel.

Aufgabe 5.11. Der Kehrwert von 1 ist 1. (Siehe auch Aufgabe 5.2.) Die Kehrwerte von anderen ganzen Zahlen sind in der Regel jedoch keine ganzen Zahlen. Es gibt allerdings eine Ausnahme, also eine Zahl $a \in \mathbb{Z} \setminus \{1\}$, für die $1/a \in \mathbb{Z}$ gilt. Welche ist es?

Inverse Elemente sind unter anderem beim Lösen von Gleichungen wichtig. Eine Gleichung wie $3x - 10 = 11$ ist eine Aussageform. Und beim *Lösen* einer solchen Gleichung geht es darum, Werte zu finden, die man für x einsetzen kann, um die Aussageform wahr zu machen. Wenn Sie „nach x auflösen“, heißt das, dass Sie die Gleichung so umformen, dass schließlich auf einer Seite des Gleichheitszeichens nur noch x steht. Dann steht nämlich auf der anderen Seite der Wert, den Sie suchen.

Ein grundlegender Gedanke beim Umformen von Gleichungen ist der folgende: Auf beiden Seiten des Gleichheitszeichens steht dasselbe; das ist ja gerade der Sinn dieses Symbols. Wenn man also mit beiden Seiten dasselbe macht, dann steht danach auf beiden Seiten immer noch dasselbe. Wir machen das nun mit $3x - 10 = 11$ und gehen dabei ganz langsam vor, um klarzustellen, was genau passiert. Zunächst ersetzen wir die Abkürzung und erhalten $3x + (-10) = 11$. Um -10 „loszuwerden“, addieren wir das zu -10 inverse Element 10 auf beiden Seiten. Das führt zu der Gleichung $3x + (-10) + 10 = 11 + 10$ bzw. $3x = 21$. Dabei gab es den entscheidenden Zwischenschritt $3x + 0 = 21$. Die Addition von 10 führte dazu, dass aus der Zahl -10 die Zahl null wurde, die wir anschließend einfach weglassen können, weil sie das neutrale Element der Addition ist.

Mit dem Rest $3 \cdot x = 21$ gehen wir nun genauso vor, wobei es nun um Multiplikation statt Addition geht. Wir wollen die Zahl 3 eliminieren und multiplizieren dafür mit ihrem Kehrwert $1/3$. Das führt zur gesuchten Lösung $x = 7$. Auch hier ist es für das Verständnis relevant, sich kurz den letzten Zwischenschritt $1 \cdot x = 7$ zu vergegenwärtigen. Die Eins spielt hier dieselbe Rolle wie die Null davor.

Aufgabe 5.12. Geben Sie die Menge $A = \{x \in \mathbb{Q} : 3x + 10 = 13 - 4x\}$ in aufzählender Schreibweise an.

Aufgabe 5.13. Geben Sie die Menge $B = \{x \in \mathbb{Q} : (3x - 4)(2x + 5) = 0\}$ in aufzählender Schreibweise an. (Hinweis: Korollar 5.2.)

Aufgabe 5.14. Geben Sie die Menge $C = \{n \in \mathbb{N} : (2n + 10)(2n - 10) = 0\}$ in aufzählender Schreibweise an.

★Aufgabe 5.15. Geben Sie die Menge $D = \{n \in \mathbb{N} : (n + 3)(n + 4) = 42\}$ in aufzählender Schreibweise an, ohne auf Ihr Schulwissen über quadratische Gleichungen zurückzugreifen.

Definition

Der **Absolutbetrag** bzw. **Betrag** (engl. *absolute value* oder *modulus*) einer rationalen Zahl x wird definiert als

$$|x| = \begin{cases} x & x \geq 0 \\ -x & x < 0 \end{cases}$$

Zum Beispiel gilt $|-5| = 5$ und $|3/4| = 3/4$. Man kann sich den Betrag vorstellen als die Funktion, die „das Vorzeichen wegnimmt“. Aber eine für die Zukunft bessere Vorstellung ist, dass der Betrag den Abstand einer Zahl von der Null auf der **Zahlengeraden** angibt. Offensichtlich ist der Betrag einer Zahl niemals negativ und null ist die einzige Zahl, deren Betrag null ist.

Was wir gerade gesehen haben, ist eine Definition durch *Fallunterscheidung*: Was $|x|$ sein soll, hängt davon ab, ob $x \geq 0$ oder $x < 0$ gilt. Das entspricht einer **bedingten Anweisung** beim Programmieren. Die einzelnen Fälle werden untereinander aufgeführt und die jeweiligen Bedingungen stehen rechts. Es kann durchaus auch mehr als zwei Fälle geben. Wichtig ist jedoch, dass sich die Bedingungen gegenseitig ausschließen und dass alle möglichen Fälle abgedeckt sind.

Fallunterscheidung

Aufgabe 5.16. Welche der folgenden Definitionen durch Fallunterscheidung (bei denen x jeweils für eine rationale Zahl steht) sind fehlerhaft und warum?

$$f(x) = \begin{cases} 3x & x \geq 3 \\ 2x+1 & x \leq 3 \end{cases} \quad g(x) = \begin{cases} 3x & x \geq 2 \\ 2x+2 & x \leq 2 \end{cases} \quad h(x) = \begin{cases} 3x & x \geq 3 \\ 2x & x \leq 2 \end{cases}$$

Für alle $x, y \in \mathbb{Q}$ gilt:

- (i) $|xy| = |x| \cdot |y|$
- (ii) $|x/y| = |x|/|y|$ für $y \neq 0$
- (iii) $|x+y| \leq |x| + |y|$
- (iv) $|x-y| \leq |x| + |y|$
- (v) $||x| - |y|| \leq |x \pm y|$

Lemma 5.7

Die beiden ersten Eigenschaften besagen, dass man Betrag und Punktrechnung vertauschen kann. Man nennt das auch die *Multiplikativität* des Betrages. Bei Eigenschaft (iii) ist besonders wichtig, dass dort anders als bei der Multiplikation *nicht* das Gleichheitszeichen steht. Das ist die sogenannte *Dreiecksungleichung*.²⁸ Die letzten beiden Eigenschaften sind Folgerungen aus dieser Ungleichung. Das Zeichen $\pm$ bedeutet, dass man an dieser Stelle sowohl $+$ als auch $-$ einsetzen kann.

Multiplikativität

Dreiecksungleichung

Aufgabe 5.17. Geben Sie Zahlen x und y an, für die $|x+y| = |x| + |y|$ *nicht* gilt.

Für $x \in \mathbb{Q}$ und $n \in \mathbb{N}^+$ ist die n -te **Potenz** (engl. *power*) x^n eine Abkürzung für

$$\underbrace{x \cdot x \cdot \dots \cdot x}_{n\text{-mal}}$$

und x^{-n} steht für den Kehrwert von x^n , wenn x nicht null ist. x^0 ist per definitionem immer 1.

Definition

²⁸Der Grund für diesen Namen wird sich **später** klären.

Basis, Exponent

Im Ausdruck a^b wird a *Basis* genannt und b *Exponent*. Beschränkt man sich zunächst auf positive Exponenten, so sind gewisse Rechenregeln wegen der **Assoziativität und Kommutativität** der Multiplikation offensichtlich. Sie werden hier am Beispiel vorgeführt:

$$2^5 \cdot 2^3 = (2 \cdot 2 \cdot 2 \cdot 2 \cdot 2) \cdot (2 \cdot 2 \cdot 2) = 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 \cdot 2 = 2^{5+3}$$

$$2^5 \cdot 3^5 = (2 \cdot 2 \cdot 2 \cdot 2 \cdot 2) \cdot (3 \cdot 3 \cdot 3 \cdot 3 \cdot 3) = 2 \cdot 3 \cdot 2 \cdot 3 \cdot 2 \cdot 3 \cdot 2 \cdot 3 \cdot 2 \cdot 3 = (2 \cdot 3)^5$$

$$(3^4)^3 = (3 \cdot 3 \cdot 3 \cdot 3)^3 = (3 \cdot 3 \cdot 3 \cdot 3) \cdot (3 \cdot 3 \cdot 3 \cdot 3) \cdot (3 \cdot 3 \cdot 3 \cdot 3) \\ = 3 \cdot 3 \cdot 3 \cdot 3 \cdot 3 \cdot 3 \cdot 3 \cdot 3 \cdot 3 \cdot 3 \cdot 3 \cdot 3 = 3^{4 \cdot 3}.$$

Wenn Sie dieses Prinzip verstanden haben, müssen Sie keine Formeln auswendig lernen, sondern Sie *sehen*, warum es so ist.

Die Definitionen für x^0 und für negative Exponenten sind nicht willkürlich gewählt, sondern sie wurden so „gebaut“, dass die obigen Regeln auch für solche Ausdrücke gelten:

Lemma 5.8

Für alle $a, b \in \mathbb{Q}$ und alle $m, n \in \mathbb{Z}$ gelten die folgenden Aussagen:

- (i) $a^m \cdot a^n = a^{m+n}$
- (ii) $a^m / a^n = a^{m-n}$ für $a \neq 0$
- (iii) $a^m \cdot b^m = (ab)^m$
- (iv) $a^m / b^m = (a/b)^m$ für $b \neq 0$
- (v) $(a^m)^n = a^{mn}$

Aufgabe 5.18. Geben Sie Zahlen a und b an, für die $a^b = b^a$ nicht gilt.

Aufgabe 5.19. Wie sehen Zehnerpotenzen aus, also Zahlen der Form 10^n für $n \in \mathbb{Z}$?

Aufgabe 5.20. Wer Informatik studiert, muss ständig mit Zweierpotenzen jonglieren. Lernen Sie in Ihrem eigenen Interesse zumindest die neun Potenzen 2^2 bis 2^{10} auswendig.

Aufgabe 5.21. Berechnen Sie ohne Taschenrechner:

- (i) $2 \cdot 3^3$
- (ii) $5^3 - 5^2$
- (iii) $(2 + 3)^3$
- (iv) $(2 \cdot 3)^3$
- (v) 4^{-3}
- (vi) $2^3 \cdot 3^2$
- (vii) $(1/3)^4$
- (viii) $3^3 \cdot 2^2$
- (ix) $7^8 / 7^{10}$
- (x) $(4/9)^2 \cdot (3/2)^4$

Aufgabe 5.22. Welche der folgenden Aussagen sind wahr?

- (i) $\forall a \in \mathbb{Q} (a^4)^7 = a^{11}$
- (ii) $\forall a \in \mathbb{Q} a^8 - a = a^7$
- (iii) $\forall a \in \mathbb{Q} (-a)^2 = -a^2$
- (iv) $\forall a \in \mathbb{Q} \setminus \{0\} a^{-5} = -a^5$

Aufgabe 5.23. Die in diesem Abschnitt angesprochenen Eigenschaften und Rechenregeln sollten Sie alle aus der Schule kennen. Die Erfahrung lehrt jedoch, dass bei vielen Studierenden gerade hier ein großes Fehlerpotential schlummert. Ich kann Ihnen daher nur raten, sich jeden einzelnen Punkt in Ruhe anzuschauen und sich zu überlegen, ob Ihnen die jeweilige Aussage wirklich klar ist oder ob Sie sie nur überflogen haben. Wenn Sie in diesem Bereich unsicher sind, dann schauen Sie *rechtzeitig* in eines der [am Ende des Skripts](#) aufgeführten Bücher zum Auffrischen des Schulstoffs und arbeiten Sie damit. Insbesondere sollten Sie sich, wenn Sie bei den Aufgaben Fehler gemacht haben, nicht einreden, dass Sie das ja „eigentlich“ können. Wenn man sich seine eigenen Defizite nicht eingesteht, kann man sie auch nicht abstellen.

Bemerkungen

Sogenannte „unechte“ Brüche werden in der Hochschulmathematik in der Regel nicht als „gemischte Zahlen“ geschrieben. Man schreibt also z. B. $\frac{5}{2}$ und *nicht* $2\frac{1}{2}$.

Da Schreibweisen wie $\frac{a}{b}$ und a/b ganz allgemein für Division verwendet werden, muss man aufpassen: Etwas, was *aussieht* wie ein Bruch, muss deshalb noch kein Bruch sein. a/b ist nur dann mit Sicherheit ein Bruch, wenn a und b ganze Zahlen sind. (Wir werden [später](#) noch sehen, dass beispielsweise $\sqrt{2}$ keine rationale Zahl ist. Setzt man $a = \sqrt{2}$ und $b = 1$, so ist a/b also *kein* Bruch.)

Der Multiplikationspunkt wird fast immer weggelassen, wenn das nicht zu Missverständnissen führt. Man schreibt also $3n$ statt $3 \cdot n$, ab statt $a \cdot b$ oder $3(x + y)$ statt $3 \cdot (x + y)$, aber natürlich *nicht* 42 statt $4 \cdot 2$.

Wie immer werden Ausdrücke in Klammern zuerst und „von innen nach außen“ ausgewertet. Für das Einsparen von Klammern gilt die Merkregel „Punkt- vor Strichrechnung“, womit gemeint ist, dass die Operationen, die man mit Punkten schreibt ($\cdot$ für Multiplikation und $:$ für Division) höhere Priorität haben als die, die man mit Strichen schreibt ($+$ für Addition und $-$ für Subtraktion). $a + bc$ steht also abkürzend für $a + (b \cdot c)$ und *nicht* für $(a + b) \cdot c$.

Bei Verknüpfungen gleicher Priorität wird implizit von links nach rechts geklammermt. $a - b + c$ steht also beispielsweise für $(a - b) + c$ und *nicht* für $a - (b + c)$. Und $a/b \cdot c$ steht für $(a/b) \cdot c$ und *nicht* für $a/(b \cdot c)$. In beiden Fällen würde die falsche Klammerung zu Rechenfehlern führen.

Divisionen mit waagerechtem (!) Strich setzen implizite Klammern um Dividend und Divisor. Ein Term wie

$$\frac{a+7}{2b-8}$$

bedeutet also dasselbe wie $(a+7)/(2b-8)$.

Potenzieren hat eine noch höhere Priorität als Punktrechnung. ab^c steht also für $a(b^c)$ und *nicht* für $(ab)^c$. Letzteren Ausdruck kann man nur *mit* Klammern schreiben. Insbesondere sind auch $(-2)^4$ und -2^4 zwei verschiedene Zahlen, nämlich 16 und -16 . Das liegt daran, dass man das Minuszeichen vor der Zwei als „Strichrechnung“ mit entsprechend geringer Priorität interpretiert.

Bei „gestapelten“ Potenzen wird anders als bei den Grundrechenarten implizit von rechts nach links geklammert. a^{b^c} steht also für $a^{(b^c)}$ und *nicht* für $(a^b)^c$.

In manchen Büchern und Fachartikeln gilt die Konvention, dass Produkte *ohne* Multiplikationspunkt eine höhere Priorität haben als „Punktrechnung“. Das würde z. B. bedeuten, dass a/bc als $a/(b \cdot c)$ gelesen wird. Das ist allerdings eine gefährliche Regel, die auch nicht allgemein üblich ist. In diesem Skript machen wir das nicht und werden im Zweifelsfall lieber „überflüssige“ Klammern setzen. So sollten Sie es auch halten.

Wenn Sie eine Variable durch einen Ausdruck ersetzen, sollten Sie diesen grundsätzlich *immer* in Klammern einpacken. (Überflüssige Klammern können Sie später noch entfernen, falls es welche gibt.) Die Formel $(a + b) \cdot c = a \cdot c + b \cdot c$ ist beispielsweise korrekt, aber wenn Sie c ohne Klammern durch $e + f$ ersetzen, erhalten Sie $(a + b) \cdot e + f = a \cdot e + f + b \cdot e + f$ und das ist falsch!

Verwenden Sie grundsätzlich Klammern, um zu vermeiden, dass zwei Rechensymbole direkt aufeinanderfolgen. Schreiben Sie also z. B. *nicht* „ $3 + -7$ “ oder „ $2 \cdot -8$ “, sondern stattdessen $3 + (-7)$ und $2 \cdot (-8)$.

Das falsche Setzen von Klammern bzw. das „Vergessen“ derselben ist einer der häufigsten Fehler in Prüfungen!

★ Arithmetik im Computer

Alle gängigen Programmiersprachen beachten die oben angesprochenen Prioritätsregeln. Geben Sie z. B. die folgenden Terme in PYTHON ein und überlegen Sie jeweils *vorher*, was nach Ihrer Meinung herauskommen sollte.

```
3+4*5
20-14//2
10-1-2
42//2*3
2*3**2
2**3**2
(2**3)**2
```

$*$ ist in der Regel der Ersatz für den Multiplikationspunkt. In PYTHON ist $**$ das Zeichen für die Potenz und $//$ das für ganzzahlige Division. Dazu [später mehr](#). Die meisten Programmiersprachen tolerieren allerdings das Weglassen des Multiplikationspunktes nicht.²⁹ Die Eingabe $6//2(1+2)$ würde in PYTHON zu einer Fehlermeldung führen.

²⁹Mir fallen als Ausnahmen nur [JULIA](#) und die [WOLFRAM LANGUAGE](#) ein.

Den Absolutbetrag gibt es in PYTHON schon als `abs`. Man könnte ihn aber auch selbst folgendermaßen programmieren, wodurch die Definition durch Fallunterscheidung deutlich wird:

```
def absval(x):
    if x>=0:
        return x
    else:
        return -x
```

Das ist allerdings auch in einer Zeile möglich:

```
absval = lambda x: x if x>=0 else -x
```

Die Potenz könnte man für ganzzahlige Exponenten **rekursiv** auch selbst programmieren, z. B. so:

```
def power(a,b):
    if b<0: return 1/power(a,-b)
    if b==0: return 1
    return a*power(a,b-1)
```

Aus didaktischen Gründen schließlich noch die rekursive Definition von Addition und Multiplikation für natürliche Zahlen, für die PYTHON nichts weiter beherrschen muss als „nächste Zahl“ und „vorherige Zahl“ (also $n + 1$ und $n - 1$):

```
add = lambda a,b: a if b==0 else 1+add(a,b-1)
mul = lambda a,b: 0 if b==0 else add(a,mul(a,b-1))
```

6. Mengen von Mengen

Wir hatten Mengen sehr allgemein als Ansammlungen von *Objekten* definiert. Mengen selbst sind auch Objekte, also können auch Mengen Elemente von Mengen sein,³⁰ z. B.

$$A = \{1, 2, \{3, 4\}\}. \quad (6.1)$$

Diese Menge besteht aus drei Objekten: den beiden Zahlen 1 und 2 sowie der Menge $\{3, 4\}$. Es gilt also insbesondere $\{1, 2\} \subseteq A$ und $\{3, 4\} \in A$. Aber es gilt *nicht* $\{1, 2\} \in A$ und auch *nicht* $3 \in A$ oder $\{3, 4\} \subseteq A$. Eigentlich sollte das allein durch das Betrachten von (6.1) klar sein, denn dort werden die drei Elemente von A ja – durch Kommata getrennt – aufgezählt. Und das dritte Element, nämlich die Menge $\{3, 4\}$ beginnt mit der sich öffnenden Mengenklammer und endet mit der

³⁰Man könnte die spitzfindige, aber völlig legitime Frage stellen, ob eine Menge ein Element von sich selbst sein kann. In der axiomatischen Mengenlehre „verbietet“ man das in der Regel durch das sogenannte *Fundierungssaxiom*. Aber es würde auch nicht schaden und für die Praxis außerhalb der Grundlagenforschung ist es nicht relevant.

zugehörigen schließenden Klammer. Aber die Erfahrung lehrt, dass es an dieser Stelle oft Verständnisprobleme gibt.

Dagegen hilft wohl nur, sich die Bedeutung von $\in$ und $\subseteq$ noch einmal klarzumachen. Und vielleicht hilft es auch, die Elemente der Menge im Geiste durch Variablen zu ersetzen. Schreibt man x und y für 1 und 2 sowie B für $\{3, 4\}$, dann ist A die Menge $\{x, y, B\}$. Jetzt *sieht* man, dass B ein Element von A ist, oder? Dass 3 kein Element von A ist (und deshalb B auch keine Teilmenge von A sein kann), liegt daran, dass 3 weder x noch y noch B ist. Sollten Sie mit dem Konzept immer noch Schwierigkeiten haben, dann lesen Sie eventuell zu schnell. Vielleicht noch einmal die [Gebrauchsanweisung](#) anschauen?

Aufgabe 6.1. Denken Sie sich zwei Mengen A und B mit $A \in B$ und $A \subseteq B$ aus.

Wir führen nun eine spezielle Menge von Mengen ein:

Definition

Die **Potenzmenge** (engl. *power set*) einer Menge A ist die Menge aller Teilmengen von A : $\mathcal{P}(A) = \{X : X \subseteq A\}$.

Wie muss man sich das vorstellen? Fangen wir ganz klein an. Wie sieht die Potenzmenge der leeren Menge aus? Die leere Menge ist – siehe Lemma 2.1 – Teilmenge *jeder* Menge, also insbesondere von sich selbst. Und andere Teilmengen kann $\emptyset$ natürlich nicht haben. Also ergibt sich $\mathcal{P}(\emptyset) = \{\emptyset\}$. Achtung! $\{\emptyset\}$ ist *nicht* dasselbe wie $\emptyset$. Die erste Menge hat genau ein Element, nämlich $\emptyset$, während die zweite *kein* Element hat. Es gilt also $\emptyset \in \{\emptyset\}$ und $\emptyset \notin \emptyset$. Darum sind die beiden Mengen nicht gleich, denn zwei Mengen sind – [Sie erinnern sich?](#) – dann und nur dann gleich, wenn sie dieselben Elemente haben.

Vielleicht kommen Ihnen solche Feinheiten pedantisch vor. Wen interessiert der Unterschied zwischen $\emptyset$ und $\{\emptyset\}$? Für einen erfolgreichen Umgang mit der Mathematik ist es jedoch wichtig, sich an eine Ausdrucksweise zu gewöhnen, die präziser als die Alltagssprache ist. Sie müssen Ihre Gedanken eindeutig und unmissverständlich artikulieren können. Und Sie müssen in der Lage sein, in mathematischen Texten nur das zu lesen, was dort wirklich steht, ohne etwas hineinzuzinterpretieren. Beides gelingt nur durch Übung.

Aber zurück zur Potenzmenge. Die nächstgrößere Menge ist eine, die nur aus einem Element besteht, z. B. $\{42\}$, also Singleton 42.³¹ Es ist wohl offensichtlich, dass es in diesem Fall zwei Teilmengen gibt, also

$$\mathcal{P}(\{42\}) = \{\emptyset, \{42\}\}.$$

Und wenn wir eine Menge wie $\{23, 42\}$ mit zwei Elementen haben, dann ergibt sich als Potenzmenge

$$\mathcal{P}(\{23, 42\}) = \{\emptyset, \{23\}, \{42\}, \{23, 42\}\}.$$

³¹Es wäre falsch, einfach „42“ zu sagen, weil die Menge $\{42\}$ nicht dasselbe wie die Zahl 42 ist. Siehe Seite 19.

Ich zeige auch noch ein Beispiel für eine dreielementige Menge und hebe dabei das Element hervor, das gegenüber dem letzten Beispiel hinzugekommen ist:

$$\mathcal{P}(\{7, 23, 42\}) = \{\emptyset, \{7\}, \{23\}, \{7, 23\}, \{42\}, \{7, 42\}, \{23, 42\}, \{7, 23, 42\}\}. \quad (6.2)$$

Aufgabe 6.2. Können Sie ein Muster erkennen? Welcher Zusammenhang besteht zwischen der Anzahl der Elemente von A und der Anzahl der Teilmengen von A ?

Vielleicht sind Sie beim Nachdenken über die letzte Aufgabe von selbst darauf gekommen: Wenn man zu einer Menge ein Element hinzufügt, dann verdoppelt sich die Anzahl der Teilmengen.³² Den Grund kann man am Beispiel (6.2) erkennen: Die Teilmengen von $\{23, 42\}$ sind auch Teilmengen von $\{7, 23, 42\}$, aber zu jeder dieser „alten“ Teilmengen kommt durch Ergänzen des Elements 7 eine „neue“ Teilmenge hinzu – aus $\emptyset$ wird $\{7\}$, aus $\{23\}$ wird $\{7, 23\}$ und so weiter.

Dieses Ergebnis wollen wir festhalten und führen dafür eine Schreibweise ein.³³

Für eine endliche Menge A bezeichnen wir mit $|A|$ die Anzahl ihrer Elemente. Man nennt das die **Mächtigkeit** (engl. *cardinality*) von A .

Definition

Was wir uns vor dieser Definition überlegt haben, können wir nun kurz und knackig so aufschreiben:³⁴

Mächtigkeit der Potenzmenge

Für jede endliche Menge A gilt $|\mathcal{P}(A)| = 2^{|A|}$.

Satz 6.1

Aufgabe 6.3. Sei A die Menge $\{\emptyset, \{\emptyset\}$ und $B = \{1, 2\}$. Geben Sie die vier Mengen $A \setminus \emptyset$, $A \setminus \{\emptyset\}$, $B \setminus \emptyset$ und $B \setminus \{\emptyset\}$ an.

Aufgabe 6.4. Welche der folgenden Aussagen ist wahr? Geben Sie bei den falschen Aussagen jeweils ein Element an, das in der einen, aber nicht in der anderen Menge enthalten ist.

- (i) $\{1, \{2\}\} = \{2, \{1\}\}$
- (ii) $\{1, \{2\}\} \cup \{2, \{1\}\} = \{1, 2\}$
- (iii) $\{1, \{2\}\} \cup \{2, \{1\}\} = \{\{1, 2\}\}$
- (iv) $\{1, \{2\}\} \cup \{2, \{1\}\} = \{\{1\}, \{2\}\}$
- (v) $\{1, \{2\}\} \setminus \{2, \{1\}\} = \{2, \{1\}\} \setminus \{1, \{2\}\}$
- (vi) $\{\emptyset\} \cup \emptyset = \{\emptyset\}$
- (vii) $\{1\} \cup \emptyset = \{1, \emptyset\}$

³²Vorsicht! Das gilt nur für endliche Mengen. Lesen Sie weiter!

³³Wir setzen eine intuitive Vorstellung davon voraus, wann eine Menge *endlich* ist. Wenn man Mengenlehre „ernsthaft“ betreibt, dann muss man erstens präzise definieren, was man mit *endlich* meint, und man kann zweitens auch für unendliche Mengen deren *Kardinalität* (ein anderes Wort für Mächtigkeit) angeben. Übrigens findet man statt $|A|$ in manchen Büchern auch die Notation $\#A$.

³⁴Das erklärt den Namen *Potenzmenge*. Wegen dieses Zusammenhangs wird in manchen Büchern auch 2^A statt $\mathcal{P}(A)$ geschrieben. Wir werden diese Schreibweise nicht verwenden.

(viii) $\{1\} \cup \{\emptyset\} = \{1\}$

(ix) $\mathcal{P}(\mathbb{N}) \setminus \{\{0\}\} = \mathcal{P}(\mathbb{N}^+)$

Aufgabe 6.5. Sei $A = [3]$ und $B = \{mn : m, n \in A\}$. Geben Sie $|B|$ an. Worauf müsste man zusätzlich achten, wenn in der Aufgabenstellung statt 3 zum Beispiel 6 stünde?

Aufgabe 6.6. Sei $A = [6]$ und $B = \{mn : m, n \in A\}$. Geben Sie die Mächtigkeit von $\mathcal{P}(\{B\})$ an.

Aufgabe 6.7. Wenn A und B Mengen mit $A \cap B = \emptyset$ sind, dann ist auch der Durchschnitt von $\mathcal{P}(A)$ und $\mathcal{P}(B)$ leer.³⁵ Stimmt das oder stimmt das nicht? Oder stimmt es nur für bestimmte Mengen?

Aufgabe 6.8. Die Aussageformen $\mathcal{P}(A \setminus B) = \mathcal{P}(A) \setminus \mathcal{P}(B)$ und $\mathcal{P}(A \cup B) = \mathcal{P}(A) \cup \mathcal{P}(B)$ sind beide nicht für beliebige Mengen A und B wahr. Finden Sie jeweils konkrete Beispiele für A und B , die sie zu falschen Aussagen machen.

Aufgabe 6.9. Seien a, b und c ganze Zahlen. Was kann man über die Mächtigkeit von $X = \{a, b, c\} \setminus \{a\}$ aussagen?

Aufgabe 6.10. Geben Sie von den folgenden Mengen jeweils die Mächtigkeit an.

$A = \emptyset$	$B = \{\emptyset\}$	$C = \{\{\emptyset\}\}$
$D = \{1, 2, 3\}$	$E = \{\{1, 2, 3\}\}$	$F = \{1, \{2, 3\}\}$
$G = \{1, \{2, \{3\}\}\}$	$H = \{\emptyset, 1, \{1\}, \{1, 1\}\}$	$I = \{\{\emptyset\}, \{2\}, \{2, \{2\}\}\}$
$J = \{1, 2\} \cup \{2, 4\}$	$K = \{1, 2\} \cap \{2, 4\}$	$L = \{1, 2\} \setminus \{2, 4\}$
$M = \{1, 2, 3\} \setminus \{2, 3\}$	$N = \{1, 2, 3\} \setminus \{\{2, 3\}\}$	$O = \{1, 2, 3\} \setminus \{\{2\}, 3\}$
$P = \{1, \{2, 3\}\} \setminus \{2, 3\}$	$Q = \{1, 2, \{3\}\} \setminus \{\{2\}, \{3\}\}$	$R = \{1, 2, \{1, 2, 3\}\} \setminus \{\{1, 2\}\}$

Aufgabe 6.11. Sortieren Sie die folgenden Mengen nach Mächtigkeit:

$A = \{n \in \mathbb{N} : n^2 = 42\}$	$B = \{n \in \mathbb{N} : n^2 + 1 < 42\}$
$C = \{n \in \mathbb{N} : n^2 - 7 = 42\}$	$D = \{n \in \mathbb{N} : (n+1)^2 < 42\}$
$E = \{n \in \mathbb{N} : n \text{ teilt } 42\}$	$F = \{n \in \mathbb{N} : 2n < 42\}$

Aufgabe 6.12. Wie viele Teilmengen von $A = \{n \in \mathbb{N} : n < 20\}$ enthalten mindestens eine Primzahl?

Aufgabe 6.13. Wie viele Teilmengen von $A = \{k \in \mathbb{N} : 3 \leq k \leq 30\}$ gibt es, die mindestens eine durch 3 teilbare Zahl enthalten?

³⁵Das ist wieder so eine typische Sprechweise, die sich eingebürgert hat. Wenn man sagt, dass eine Menge *leer* ist, meint man damit natürlich, dass es sich um *die* leere Menge handelt.

★ Mengen von Mengen im Computer

PYTHON bietet zusätzlich zu Mengen „gefrorene Mengen“ an. Damit kann man Mengen als Elemente von Mengen speichern, was mit dem Datentyp `set` nicht möglich ist.³⁶

```
X = frozenset({1,2,3})
Y = {1,2,3,X}
X <= Y, X in Y
```

Eine Aufgabe wie 6.5 würde man für größere Zahlen wohl lieber bequem mit Computerhilfe lösen. Das könnte so aussehen:

```
A = set(range(1,7))
len({m*n for m in A for n in A})    # Mächtigkeit einer Menge
```

Die folgende Funktion ist ein **Generator**, der **rekursiv** durch die Potenzmenge einer Menge iteriert.³⁷

```
def powerSet(M):                                # engl. Wort für Potenzmenge
    if M == set():
        yield set()
    else:
        S = {next(iter(M))}                    # {x} für ein x aus M
        for A in powerSet(M - S):
            yield A
            yield A | S
```

Dass Satz 6.1 zumindest für die ersten zehn natürlichen Zahlen stimmt, kann man so überprüfen:

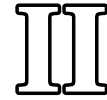
```
all(2**n == sum(1 for A in powerSet(set(range(n))))
    for n in range(10))
```

Will man wirklich die Potenzmenge haben, dann muss man auf die eben erwähnten gefrorenen Mengen zurückgreifen:

```
{frozenset(X) for X in powerSet({1,23,42})}
```

³⁶Wenn Sie in dem folgenden Beispiel `frozenset({1,2,3})` durch `{1,2,3}` ersetzen, erhalten Sie eine Fehlermeldung. In Mengen kann man nur Objekte speichern, die nicht *mutable* (siehe Seite 25) sind, und gefrorene Mengen sind solche Objekte: `add` und `remove` erlauben sie nicht. In diesem Sinne sind sie „mathematischer“ als die normalen Mengen von PYTHON. (Für die PYTHON-Experten: Eigentlich geht es darum, ob die Objekte *hashable* sind.)

³⁷In der kommentierten Zeile wird ein beliebiges Element von `M` ausgewählt. Die Konstruktion ist etwas umständlich, da man nicht einfach so etwas wie `M[0]` schreiben kann. (Weil Mengen auch in PYTHON keine festgelegte Reihenfolge haben.)



Zahlentheorie

Wir werden uns nun mit den Grundzügen der sogenannten *Zahlentheorie* befassen. Abgesehen davon, dass es sich einfach um ein spannendes Thema handelt, ist die Zahlentheorie auch für die Anwendung relevant, weil sie eine zentrale Rolle in der **Kryptographie** spielt. Zudem kann man einige ihrer Techniken wie z. B. die **modulare Arithmetik** gewinnbringend in der Informatik einsetzen.

Das Wort *Zahlentheorie* (engl. *number theory*) stammt aus einer Zeit, als man nur als Zahlen ansah, was wir heutzutage als positive natürlichen Zahlen bezeichnen (während man Brüche oder irrationale Zahlen *Größen* nannte). Wir werden das etwas ausweiten und teilweise auch mit negativen Zahlen arbeiten, aber andere Zahlen werden auf den folgenden Seiten keine Rolle spielen.

Achtung: Im gesamten Kapitel geht es ausschließlich um ganze Zahlen.¹ Wenn eine Variable für eine Zahl steht und nichts weiter über sie gesagt wird, dann wird implizit immer vorausgesetzt, dass sie ein Element von $\mathbb{Z}$ ist. Und wenn von *Zahlen* ohne weiteren Zusatz die Rede ist, dann sind damit ganze Zahlen gemeint.

7. Teilbarkeit

Wir wissen inzwischen, dass ganze Zahlen in der Regel keinen Kehrwert haben. Man kann also nicht dividieren, wenn man die Menge $\mathbb{Z}$ nicht verlassen will. Es gibt allerdings eine ähnliche Operation, die man typischerweise bereits in der Grundschule lernt: das Teilen mit Rest.

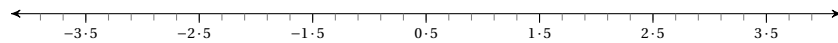
Division mit Rest

Ist $m \in \mathbb{Z} \setminus \{0\}$, so gibt es zu jedem $a \in \mathbb{Z}$ eindeutig bestimmte Zahlen $k \in \mathbb{Z}$ und $r \in \mathbb{N}$ mit $r < |m|$ und $a = km + r$.

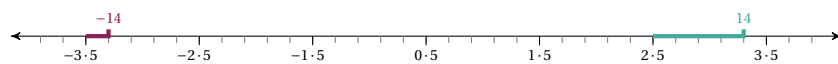
Satz 7.1

¹Eine Ausnahme bilden lediglich die in Abschnitt 9 eingeführten Abkürzungen für Restklassen, die „aussehen“ wie Zahlen. Den Unterschied werden Sie dann aber hoffentlich erkennen.

Dass das stimmt, wird grafisch recht einfach deutlich. Für $m = 5$ würde man beispielsweise an allen Vielfachen von 5 auf der Zahlengeraden Markierungen anbringen.



Dadurch bilden sich (unendlich viele) „Schubfächer“ der Breite m , zum Beispiel das von $1 \cdot 5$ bis (vor) $2 \cdot 5$, das die Zahlen 5, 6, 7, 8 und 9 enthält. Die Zahl a muss in einem dieser Schubfächer landen und ist damit auf jeden Fall höchstens $m - 1$ Einheiten von dessen linkem Rand entfernt. Für die Fälle $a = 14$ und $a = -14$ sieht es so aus:



Definition

Die Zahl r aus Satz 7.1 wird als **Rest** bei der Division von a durch m bezeichnet und $a \bmod m$ geschrieben sowie als „ a modulo m “ gelesen.

Wir haben also:

$$\begin{array}{ll} 14 \bmod 5 = 4 & 14 = 2 \cdot 5 + 4 \\ -14 \bmod 5 = 1 & -14 = -3 \cdot 5 + 1 \\ 14 \bmod -5 = 4 & 14 = -2 \cdot (-5) + 4 \\ -14 \bmod -5 = 1 & -14 = 3 \cdot (-5) + 1 \end{array}$$

Die ersten beiden Zeilen kann man der obigen Grafik entnehmen. Die anderen beiden basieren auf unserer Definition, dass $a \bmod m$ nie negativ ist. Da besteht jedoch leider mal wieder keine Einigkeit. Zum Glück ist zumindest der Gebrauch für positive m in der Mathematik einheitlich und nur für solche werden wir das auch verwenden.² Wichtig ist jedenfalls, dass $-14 \bmod 5$ sich *nicht* einfach dadurch ergibt, dass man ein Minuszeichen vor $4 = 14 \bmod 5$ setzt!

Aufgabe 7.1. Rechnen Sie ein paar Beispiele durch, um sich mit dieser Operation vertraut zu machen. Sie können Ihre Ergebnisse unter <https://weitz.de/mod.html> überprüfen.

Aufgabe 7.2. Wenn Sie nur einen einfachen Taschenrechner (wie [den Klausurrechner des Departments](#)) zur Verfügung haben, wird dieser höchstwahrscheinlich keine Taste für den Rest haben. Beherrscht er jedoch das Rechnen mit Brüchen, so kann man trotzdem Reste berechnen lassen. Wie?

Die Schreibweise $a \mid b$ für „ a teilt b “ hatten wir bereits eingeführt. Nun definieren wir „nachträglich“, was genau damit eigentlich gemeint sein soll. Neu ist dabei gegenüber der Schule gegebenenfalls, dass wir Teilbarkeit auch für negative Zahlen definieren.

²In vielen Programmiersprachen wird für $\bmod$ das Zeichen $\%$ verwendet. Auch dort ist man sich nicht einig. Details dazu findet man [hier](#).

Für $a, b \in \mathbb{Z}$ bezeichnet man a als **Teiler** von b bzw. b als **Vielfaches** von a , wenn es ein $k \in \mathbb{Z}$ mit $b = ka$ gibt. Man sagt auch, dass a die Zahl b **teilt**. Die Zahlen $1, -1, b$ und $-b$ nennt man **triviale Teiler** von b , alle anderen Teiler von b werden **echte Teiler** genannt.

Definition

Wie die Definition schon erahnen lässt, sind die Zahlen $1, -1, a$ und $-a$ immer Teiler von a . Das folgt aus den Gleichungen $a = 1 \cdot a$ bzw. $a = (-1) \cdot (-a)$, wobei jeder der Faktoren die Rolle von k in der Definition spielen kann. Außerdem sind zwei weitere Punkte offensichtlich:

- Weil die Aussagen $b = ka$, $b = (-k)(-a)$, $-b = (-k)a$ und $-b = k(-a)$ alle äquivalent sind, kann man sich bei Teilbarkeitsfragen auf die nichtnegativen Zahlen konzentrieren.
Konkretes Beispiel: Wenn man die positiven Teiler von 35 – also 1, 5, 7 und 35 – kennt, dann kennt man auch die negativen. Man erhält sie, indem man zu den positiven Teilern jeweils die negativen Gegenstücke hinzunimmt: $-1, -5, -7$ und -35 . Und dann kennt man auch die Teiler von -35 : es sind dieselben wie die von 35.
- Für $a \neq 0$ gilt $a \mid b$ genau dann, wenn $b \bmod a$ null ist, wenn man also b ohne Rest durch a teilen kann.

Aufgabe 7.3. Stimmt es, dass jede Zahl mindestens vier verschiedene Teiler hat?

Aufgabe 7.4. Ist 7 ein Teiler von 0? Ist 0 ein Teiler von 7?

Aufgabe 7.5. Was sind gerade bzw. ungerade Zahlen? Das wissen Sie doch, oder?

gerade ungerade

Es ist sinnvoll, ein paar einfache Teilbarkeitsregeln im Kopf zu haben:³

- Weil wir das **Dezimalsystem** verwenden, kann man eine durch 10 teilbare Zahl daran erkennen, ob die letzte Ziffer eine Null ist.
- Weil 2 und 5 Teiler von 10 sind, erkennt man Teilbarkeit durch 2 bzw. 5 daran, ob die letzte Ziffer durch 2 bzw. 5 teilbar ist.
- Weil 4 ein Teiler von 100 ist, erkennt man Teilbarkeit durch 4 daran, ob die letzten beiden Ziffern eine durch 4 teilbare Zahl darstellen: 7324 ist z. B. durch 4 teilbar, weil 24 durch 4 teilbar ist.
- Ebenso ist eine Zahl genau dann durch 8 teilbar, wenn die letzten drei Ziffern eine durch 8 teilbare Zahl darstellen.
- Eine Zahl ist genau dann durch 3 bzw. 9 teilbar, wenn ihre (iterierte) Quersumme durch 3 bzw. 9 teilbar ist. Warum das so ist, werden wir **später** klären. Die (dezimale) **Quersumme** einer natürlichen Zahl ist die Summe ihrer Ziffern (im Dezimalsystem), also ist z. B. $18 = 9 + 4 + 5$ die Quersumme von 945. Die *einstellige* oder *iterierte Quersumme* einer Zahl erhält man, indem man so lange die Quersumme bildet, bis das Ergebnis einstellig ist. Die iterierte

Quersumme

³Siehe dazu auch Aufgabe 7.7.

Quersumme von 945 ist also 9, die Quersumme von 18. (Daraus folgt, dass 945 durch 3 und durch 9 teilbar ist.)

- Eine Zahl ist durch 6 teilbar, wenn sie durch 2 *und* durch 3 teilbar ist.
- Auch für die Zahlen 7 und 11 gibt es Teilbarkeitsregeln. Die sind aber etwas komplizierter und werden ebenfalls erst *später* besprochen.

Siehe dazu auch das Vorkurs-Video *Teilbarkeit, Teilen mit Rest*.

Aufgabe 7.6. Welche der folgenden Aussagen sind wahr?

- (i) 300357 ist durch 3 teilbar.
- (ii) 300357 ist durch 6 teilbar.
- (iii) 4 teilt 20412342.
- (iv) 5 teilt 987612345.
- (v) 399168 ist ein Vielfaches von 8.
- (vi) 399168 ist ein Vielfaches von 9.

Es folgen zwei grundlegende Eigenschaften der Teilbarkeitsrelation. Erstens die sogenannte *Transitivität*:

Lemma 7.2

Aus $a \mid b$ und $b \mid c$ folgt $a \mid c$.

Warum das stimmt, kann man bereits an einem Beispiel ganz gut erkennen. 7 teilt 21, und 21 teilt 84, darum muss nach dem Lemma auch 7 ein Teiler von 84 sein. 7 teilt 21, weil $21 = 3 \cdot 7$ gilt, und 21 teilt 84, weil $84 = 4 \cdot 21$ gilt. Setzt man beides zusammen, so erhält man

$$84 = 4 \cdot 21 = 4 \cdot (3 \cdot 7) = (4 \cdot 3) \cdot 7 = 12 \cdot 7$$

und hat die Begründung dafür, dass 84 ein Vielfaches von 7 ist.

Lemma 7.3

Wenn a sowohl b als auch c teilt, dann teilt a auch $xb + yc$, wobei x und y beliebige (ganze) Zahlen sind.

Auch das machen wir uns mit einem Beispiel klar. 7 teilt sowohl 21 als auch 35. Wenn das Lemma stimmt, dann muss 7 unter anderem auch $20 \cdot 21 - 5 \cdot 35 = 245$ (mit $x = 20$ und $y = -5$) teilen. Und das stimmt auch:

$$\begin{aligned} 245 &= 20 \cdot 21 - 5 \cdot 35 = 20 \cdot (3 \cdot 7) - 5 \cdot (5 \cdot 7) = (20 \cdot 3) \cdot 7 - (5 \cdot 5) \cdot 7 \\ &= (20 \cdot 3 - 5 \cdot 5) \cdot 7 = 35 \cdot 7 \end{aligned}$$

Dabei hätten wir den Wert 35 gar nicht ausrechnen müssen, denn schon im Schritt davor war klar, dass 7 ein Teiler von 245 ist.

Linearkombination

Einen Ausdruck der Form $xb + yc$ nennt man eine *Linearkombination* von b und c .⁴ Beachten Sie, dass u. a. auch Fälle wie $x = y = 1$ oder $x = 1$ und $y = -1$ durch

⁴Genauer: eine Linearkombination mit ganzzahligen Koeffizienten.

Lemma 7.3 abgedeckt werden. Die Summe $b + c$ und die Differenz $b - c$ sind also insbesondere auch durch a teilbar, wenn b und c es sind.

Aufgabe 7.7. Nach den Teilbarkeitsregeln von Seite 57 ist 230 durch 10 teilbar, 235 durch 5 teilbar und 236 durch 2 und 4 teilbar. Begründen Sie das etwas ausführlicher mithilfe der Lemmata 7.2 und 7.3.

Wenn c Teiler von a ist, folgt aus der Definition offenbar $|c| \leq |a|$, wenn a nicht null ist. Zwei beliebige Zahlen a und b haben also erstens immer mindestens einen *gemeinsamen* positiven Teiler (nämlich 1) und zweitens können ihre gemeinsamen Teiler nicht größer als der kleinere der beiden Werte $|a|$ und $|b|$ sein.⁵ Daher ergibt die folgende Definition Sinn.

Der **größte gemeinsame Teiler** $\text{ggT}(a, b)$ von a und b ist die größte Zahl, die a und b teilt, wenn $\{a, b\} \neq \{0\}$ gilt. Zur Vervollständigung der Definition setzen wir außerdem $\text{ggT}(0, 0) = 0$. Der größte gemeinsame Teiler von mehr als zwei Zahlen wird entsprechend definiert.

Definition

Beispielsweise hat 20 die positiven Teiler 1, 2, 4, 5, 10 und 20 und 35 hat die positiven Teiler 1, 5, 7 und 35. Die größte Zahl, die in beiden Aufzählungen vorkommt, ist 5. Also gilt $\text{ggT}(20, 35) = 5$. Offenbar ist der größte gemeinsame Teiler zweier Zahlen nie negativ und nur dann null, wenn beide Zahlen verschwinden.⁶

Aufgabe 7.8. Welchen Wert haben $\text{ggT}(-20, 35)$, $\text{ggT}(20, -35)$ und $\text{ggT}(-20, -35)$?

Aufgabe 7.9. Und was ist mit $\text{ggT}(42, 42)$ und $\text{ggT}(0, 42)$?

Wie berechnet man den größten gemeinsamen Teiler? Man könnte erst alle Teiler der beiden Zahlen ermitteln und dann den größten herauspicken, der sich in beiden Mengen befindet. Das wäre aber sehr ineffizient. Zum Glück kennt man seit der Zeit von **Euklid** ein cleveres Verfahren, das erstaunlicherweise auch über zwei Jahrtausende später noch immer der beste bekannte **Algorithmus** für diese Aufgabe ist und das in dieser Form auch von Computern eingesetzt wird. Es ist der *euklidische Algorithmus* zur Berechnung von $\text{ggT}(a, b)$.⁷

euklidischer
Algorithmus

- (i) Setze A und B auf die Werte a und b .
- (ii) Gilt $A = 0$, dann ist das Ergebnis B und der Algorithmus ist beendet.
- (iii) Ist $B = 0$, dann ist das Ergebnis A und der Algorithmus ist beendet.
- (iv) Ist $A > B$, dann ersetze A durch $A - B$, anderenfalls ersetze B durch $B - A$.
- (v) Fahre fort mit Schritt (iii).

⁵Auch hier wieder mit der Ausnahme, dass nicht $a = b = 0$ gelten darf.

⁶*Verschwinden?* Was soll das denn bedeuten? Das ist mal wieder Mathematikerschnack. Wenn man sagt, dass etwas verschwindet, dann meint man damit, dass es null ist.

⁷Falls Sie sich an Seite 25 erinnern, auf der man lesen kann, dass sich in der Mathematik die Werte von Variablen nie ändern: A und B , die sich absichtlich typographisch von a und b unterscheiden, sind keine Variablen, sondern so etwas wie Speicherplätze in einem Computer.

Schauen wir uns zunächst ein Beispiel an. Hier wird in jeder Zeile der Stand eines Durchlaufs am Punkt (iii) gezeigt.

A	B
400	225
$175 = 400 - 225$	225
175	$50 = 225 - 175$
$125 = 175 - 50$	50
$75 = 125 - 50$	50
$25 = 75 - 50$	50
25	$25 = 50 - 25$
25	$0 = 25 - 25$

Am Ende haben wir $\text{ggT}(400, 225) = 25$ errechnet.

Aufgabe 7.10. Rechnen Sie selbst ein paar Beispiele durch. Ruhig auch mit größeren Zahlen. Der Algorithmus wird trotzdem immer nach recht wenigen Schritten zu einem Ergebnis kommen. Sie können Ihre Resultate überprüfen, indem Sie z. B. in [Wolfram Alpha](#) so etwas wie $\text{gcd}(400, 225)$ eingeben.⁸

Wenn man das Verfahren analysiert, dann kann man sich überzeugen, dass es bei jeder Eingabe terminiert und immer das gewünschte Ergebnis liefert. Das halten wir fest:

Satz 7.4

Der euklidische Algorithmus liefert für jede Eingabe $a, b \geq 0$ immer den größten gemeinsamen Teiler von a und b .

Diese Analyse liefert zudem eine weitere Charakterisierung des größten gemeinsamen Teilers.

Korollar 7.5

Es gilt genau dann $d = \text{ggT}(a, b)$, wenn d Teiler von a und b ist und jeder gemeinsame Teiler von a und b auch d teilt.

Beispiel: Der größte gemeinsame Teiler von $a = 79781$ und $b = 116909$ ist 221. Weitere gemeinsame Teiler der beiden Zahlen sind 1, 13 und 17. Alle drei sind auch Teiler von 221 (denn 221 ist das Produkt von 13 und 17).

Das hat zur Konsequenz, dass man den größten gemeinsamen Teiler von drei oder mehr Zahlen schrittweise berechnen kann: Um $d = \text{ggT}(79781, 116909, 215441)$ zu ermitteln, kann man beispielsweise auf das Ergebnis $d_1 = 221$ von eben zurückgreifen. Es muss $d = \text{ggT}(221, 215441) = 17$ gelten, denn erstens ist d als Teiler von d_1 auch Teiler von a und b und zweitens ist jeder gemeinsame Teiler k von a , b und $c = 215441$ insbesondere gemeinsamer Teiler von a und b und damit nach Korollar 7.5 Teiler von d_1 . So ein k muss also als Teiler von d_1 und c (erneut nach

⁸Im Englischen (und vielen Programmiersprachen) heißt es *greatest common divisor*.

Korollar 7.5) Teiler von d sein. Daher ist d offensichtlich der größte gemeinsame Teiler von a , b und c .

Aufgabe 7.11. Berechnen Sie $\text{ggT}(142\,766, 391\,989, 430\,882)$.

Wir wissen nach Lemma 7.3 bereits, dass jeder gemeinsame Teiler von a und b jede Linearkombination von a und b teilt. Aus dem euklidischen Algorithmus ergibt sich zusätzlich noch das Folgende:

Lemma von Bézout

$\text{ggT}(a, b)$ lässt sich als Linearkombination von a und b darstellen.

Korollar 7.6

Die Begründung für diese Aussage ist, dass sowohl A als auch B in jedem Schritt Linearkombinationen von a und b sind. Denn am Anfang sind sie es⁹ und in jedem Schritt danach werden Linearkombinationen der vorherigen Werte gebildet, die wieder Linearkombinationen der Ausgangswerte sind:

$$x(x_1a + y_1b) + y(x_2a + y_2b) = (xx_1 + yy_2)a + (xy_1 + yy_2)b$$

Darum ist auch der letzte Wert, also das Ergebnis des Algorithmus, eine solche Linearkombination.

Hier sieht man das noch einmal am Beispiel von vorhin:

A	B
$400 = a$	$225 = b$
$175 = a - b$	$225 = b$
$175 = a - b$	$50 = b - (a - b) = 2b - a$
$125 = (a - b) - (2b - a) = 2a - 3b$	$50 = 2b - a$
$75 = (2a - 3b) - (2b - a) = 3a - 5b$	$50 = 2b - a$
$25 = (3a - 5b) - (2b - a) = 4a - 7b$	$50 = 2b - a$
$25 = 4a - 7b$	$25 = (2b - a) - (4a - 7b) = -5a + 9b$
$25 = -5a + 9b$	0

Aufgabe 7.12. Rechnen Sie für die Werte aus Aufgabe 7.10 auch solche Linearkombinationen aus. Ob die Ergebnisse richtig sind, lässt sich ja leicht überprüfen.

★**Aufgabe 7.13.** Ist die Linearkombination, die der euklidische Algorithmus liefert, eindeutig bestimmt? (Hinweis: Wenn man $3 \cdot 14$ Schritte vorwärts und dann $14 \cdot 3$ Schritte rückwärts macht, ist man wieder am Ausgangspunkt.)

Aufgabe 7.14. Im Beispiel mit 400 und 225 wurde dreimal nacheinander 50 subtrahiert. Das hätte man auch einfacher haben können, oder? Mit anderen Worten: Man hätte doch gleich sehen können, dass 50 in 175 dreimal „reinpasst“ und dass es eigentlich um etwas anderes geht. Formulieren Sie eine verbesserte Form des euklidischen Algorithmus, die das berücksichtigt.

⁹A ist $1 \cdot a + 0 \cdot b$, B ist $0 \cdot a + 1 \cdot b$.

Definition

Zwei Zahlen werden **teilerfremd** genannt, wenn ihr größter gemeinsamer Teiler die Zahl 1 ist.

Zwei teilerfremde Zahlen haben also gewissermaßen nur die gemeinsamen Teiler, die sie haben „müssen“. Zwei verschiedene **Primzahlen** wie 7 und 13 sind immer teilerfremd. Aber auch andere Paare von Zahlen – z. B. 14 und 45 – können teilerfremd sein.

Aufgabe 7.15. Sei n irgendeine positive natürliche Zahl. Können Sie begründen, warum n und $n + 1$ garantiert teilerfremd sind?

Aufgabe 7.16. Was kann man über einen Bruch a/b sagen, wenn sein Zähler a und sein Nenner b teilerfremd sind?

Aufgabe 7.17. Nach dem **Lemma von Bézout** gibt es zu teilerfremden Zahlen a und b Zahlen x und y mit $1 = xa + yb$. Begründen Sie die umgekehrte Aussage: Wenn es Zahlen x und y mit $1 = xa + yb$ gibt, dann sind a und b teilerfremd. (Hinweis: Lemma 7.3.)

★**Aufgabe 7.18.** Kann man natürliche Zahlen x und y finden, für die $x!y! = x! + y! + 2$ gilt? (Mit dem Ausrufezeichen ist die **Fakultät** gemeint.) Begründen Sie Ihre Antwort.

★**Aufgabe 7.19.** Eine Knobelaufgabe, mit der man viele Teilbarkeitsregeln üben kann: Finden Sie eine neunstellige Zahl mit der folgenden Eigenschaft: Die aus den ersten beiden Ziffern gebildete Zahl ist durch 2 teilbar, die aus den ersten drei Ziffern gebildete Zahl ist durch 3 teilbar, die aus den ersten vier Ziffern gebildete Zahl ist durch 4 teilbar und so weiter – schließlich soll die gesamte Zahl durch 9 teilbar sein. Außerdem soll jede der Ziffern 1 bis 9 in der Dezimaldarstellung der Zahl genau einmal vorkommen.

Bemerkungen

Es ist hoffentlich klar geworden, dass es keinen Sinn ergäbe, die Aussagen aus diesem Kapitel auf alle rationalen Zahlen zu übertragen. Weil man dort immer teilen kann (außer durch null), gäbe es nie einen „Rest“. Daher wäre auch jede Zahl durch jede „teilbar“. Das wäre langweilig.

Bei den Teilbarkeitsregeln ist zu unterscheiden zwischen den Eigenschaften der Zahl und der Art und Weise, wie wir sie hinschreiben. Dass man die Zahl 72 durch 9 teilen kann, kann man an ihrer **Quersumme** erkennen. Das liegt daran, dass wir das **Dezimalsystem** verwenden. (Mehr dazu in Aufgabe 9.13.) Man kann statt 72 aber auch 1001000_2 im **Dualsystem** oder LXXII in **römischer Zahlschrift** schreiben. Dann sieht man die Teilbarkeit nicht mehr sofort, aber es ist immer noch dieselbe Zahl und sie ist immer noch durch 9 (bzw. 1001_2 bzw. IX) teilbar!

★ Teilbarkeit im Computer

In PYTHON wird das Zeichen % für mod verwendet. Die Zahl k aus Satz 7.1 kann man mit $a//m$ berechnen. In beiden Fällen wird man jedoch nicht immer dieselben

Ergebnisse wie in diesem Kapitel erhalten, wenn negative Zahlen ins Spiel kommen. Da sollten Sie vorsichtshalber die [Dokumentation](#) lesen. Den Test auf Teilbarkeit könnte man jedenfalls so programmieren:

```
divides = lambda a,b: b==0 or a!=0 and b%a==0
```

Die wortwörtliche Übersetzung des euklidischen Algorithmus von oben sieht folgendermaßen aus:

```
def ggT(a,b):  
    if a==0:  
        return b  
    while b!=0:  
        if a>b:  
            a = a-b  
        else:  
            b = b-a  
    return a
```

Aber man muss das nicht selbst machen:

```
from math import gcd  
gcd(400,225)
```

Sogar für die Faktoren aus dem [Lemma von Bézout](#) kann man sich eine fertige Funktion laden, wenn man die Bibliothek [SYMPY](#) installiert hat.¹⁰

```
from sympy import gcdex  
gcdex(400,225)
```

Und in Kapitel 5 von [Konkrete Mathematik \(nicht nur\) für Informatiker](#) sowie in den zugehörigen Lösungshinweisen findet man diverse Beispiele dafür, wie man das auch selbst programmieren kann.

8. Primzahlen

Nun kommen wir zu den Hauptakteuren der Zahlentheorie, den Primzahlen. Wir hatten [bereits definiert](#), welche Zahlen so bezeichnet werden.

Aufgabe 8.1. Es ist ganz hilfreich, wenn man zumindest die ersten zehn Primzahlen auswendig im Kopf hat: 2, 3, 5, 7, 11, 13, 17, 19, 23 und 29. Vielleicht prägen Sie sich die in einer freien Minute mal ein?

¹⁰Das Kürzel *ex* steht für *extended*. Wenn man nicht nur den größten gemeinsamen Teiler, sondern nebenbei noch diese Faktoren ausrechnet, spricht man auch vom *erweiterten* euklidischen Algorithmus.

Definition

Ein Teiler einer Zahl a , der eine Primzahl ist, wird **Primteiler** oder **Primfaktor** von a genannt.

42 hat also z. B. die Primteiler 2, 3 und 7, aber auch noch die Teiler 1, 6, 14, 21 und 42, die alle *keine* Primteiler sind. Die negativen Teiler sind natürlich „erst recht“ keine Primteiler.

Primzahlen sind ihre eigenen Primteiler. Was ist mit den zusammengesetzten Zahlen? Schauen wir uns als Beispiel $a = 21\,301\,313\,713$ an. Wenn a keine Primzahl ist, dann muss es einen Teiler b von a geben, der größer als 1 und kleiner als a ist. Durch geduldiges Probieren oder mithilfe eines Computers könnten wir herausfinden, dass $b = 2013547$ ein Teiler von a . Ist b eine Primzahl, dann ist b also ein Primteiler von a . Ist b jedoch *keine* Primzahl, dann muss es einen Teiler c von b zwischen 1 und b geben. Durch erneutes Probieren kommen wir eventuell auf $c = 19549$. Das ist immer noch keine Primzahl, aber wenn wir nicht aufgeben, dann finden wir den Teiler 173 von c und *das* ist eine Primzahl. Was wir erreicht haben, sieht so aus:

$$173 \mid 19549, \quad 19549 \mid 2013547 \quad \text{und} \quad 2013547 \mid 21\,301\,313\,713.$$

Wegen Lemma 7.2 ist klar, dass die Primzahl 173 ein Teiler von a ist. Bei einer anderen Zahl würden wir vielleicht deutlich mehr als drei Schritte brauchen, aber ein Prozess wie dieser muss offenbar jedes Mal terminieren, weil die Zahlen in jedem Schritt kleiner werden und es daher nicht immer weitergehen kann. Aufhören kann es jedoch nur mit einer Primzahl. Daher gilt die folgende Aussage:

Lemma 8.1

Jede Zahl, die größer als 1 ist, hat (mindestens) einen Primteiler.

Auf dieser grundlegenden Erkenntnis beruht der folgende Satz, der schon vor Beginn der Zeitrechnung bewiesen wurde.

Satz 8.2**Satz des Euklid**

Es gibt unendlich viele Primzahlen, d. h., $\mathbb{P}$ ist eine unendliche Menge.

Weil der Beweis dieses Satzes so schön ist, wollen wir ihn uns anschauen. Bildet man ein Produkt mehrerer Zahlen – z. B. das von a , b , c und d – und addiert man 1 dazu, so ergibt sich offenbar der Rest 1, wenn man diese Zahl durch a oder b oder c oder d teilt. Keine der Zahlen a bis d ist also ein Teiler von $1 + abcd$. Sind nun p_1 bis p_n die ersten n Primzahlen, so ist mit derselben Begründung keine von ihnen ein Teiler von $q = 1 + p_1 p_2 \cdots p_n$. Andererseits muss q nach Lemma 8.1 einen Primteiler haben. Also gibt es außer p_1 bis p_n noch mindestens eine weitere Primzahl. Und das war's schon. Wenn Sie „nur“ endlich viele Primzahlen haben, dann haben Sie garantiert nicht alle ...

Aufgabe 8.2. Ein häufiges Missverständnis ist, dass die im Beweis von Satz 8.2 definierte Zahl q selbst eine Primzahl sei. Und Zahlen wie

$$1 + p_1 = 1 + 2 = 3, \quad 1 + p_1 p_2 = 1 + 2 \cdot 3 = 7 \quad \text{und} \quad 1 + p_1 p_2 p_3 = 1 + 2 \cdot 3 \cdot 5 = 31$$

sind in der Tat Primzahlen. Begründen Sie mit einem konkreten Gegenbeispiel, dass diese Annahme trotzdem falsch ist.

Aufgabe 8.3. Begründen Sie, warum $\mathbb{N} \setminus \mathbb{P}$ ebenfalls unendlich ist.

Die folgende Aussage erklärt, warum die Primzahlen so wichtig sind. Sie sind gleichsam die „Bausteine“ der Zahlenwelt, aus denen die anderen Zahlen zusammengesetzt sind. Man könnte fast von den „Atomen der Arithmetik“ sprechen: Ein Kohlenstoffmonoxid-Molekül besteht beispielsweise immer aus einem Kohlenstoff- und einem Sauerstoffatom. Da wird nicht auf einmal ein zweites Sauerstoffatom auftauchen oder der Kohlenstoff durch Natrium ersetzt werden.¹¹ So ist das auch mit den Primzahlen.

Fundamentalsatz der Arithmetik

Jede positive ganze Zahl lässt sich eindeutig (bis auf die Reihenfolge) als Produkt von Primzahlen darstellen.

Satz 8.3

Die Zahl 2076 lässt sich beispielsweise als $2 \cdot 2 \cdot 3 \cdot 173$ darstellen. Dass die Reihenfolge nicht eindeutig ist, ist wegen der Kommutativität der Multiplikation klar: $2 \cdot 173 \cdot 3 \cdot 2$ wäre auch eine Darstellung von 2076 als Produkt von Primzahlen. Aber anders geht es nicht. Man kann keinen Faktor weglassen und man kann dasselbe Ergebnis auch nicht mit *anderen* Faktoren erhalten.

Eventuell haben Sie sich gefragt, wie man eine Primzahl oder die Zahl eins als Produkt von Primzahlen darstellen soll. Für die Praxis ist das zwar irrelevant, aber hier ist die Antwort, die Ihnen eine Mathematikerin geben würde: Der Satz sagt nicht, wie viele Faktoren das Produkt hat. Eine Primzahl lässt sich als ein Produkt darstellen, das aus *einem* Faktor besteht, und die Eins ist das **Produkt von null Primzahlen**.

Die Reihenfolge spielt zwar für das Ergebnis keine Rolle, aber man kann die an einem Produkt beteiligten Faktoren natürlich trotzdem sortieren und gleiche Faktoren zu Potenzen zusammenfassen. Damit erhält man dann die sogenannte *kanonische Primfaktorzerlegung*, für die wir hier Beispiele sehen:

kanonische
Primfaktorzerlegung

$$16 = 2^4$$

$$27 = 3^3$$

$$42 = 2 \cdot 3 \cdot 7$$

$$43 = 43$$

$$96 = 2^5 \cdot 3$$

$$1\,320\,574\,363 = 29^2 \cdot 31 \cdot 37^3$$

¹¹Allerdings hinkt der Vergleich etwas, weil es nur endlich viele **chemische Elemente** gibt, während wir gerade bewiesen haben, dass es unendlich viele Primzahlen gibt.

Aufgabe 8.4. Berechnen Sie zur Übung ein paar Primfaktorzerlegungen. Sie können Ihre Ergebnisse z. B. kontrollieren, indem Sie so etwas wie `factor 1320574363` in [Wolfram Alpha](#) eingeben.

Aufgabe 8.5. Wenn Sie nur einen einfachen Taschenrechner (wie [den des Departments](#)) zur Verfügung haben und die Primfaktorzerlegung einer großen Zahl wie 32 892 berechnen sollen, wie gehen Sie dann „strategisch“ am besten vor?

Aufgabe 8.6. In Aufgabe 8.5 kommt die Zahl 2 741 vor, von der Sie sicher nicht aus dem Kopf wissen, dass es eine Primzahl ist. Wie kann man sich mit möglichst wenigen Divisionen sicher sein, dass eine Zahl prim ist?

Welcher Zusammenhang besteht zwischen dem größten gemeinsamen Teiler zweier Zahlen und deren Primfaktorzerlegungen? Nennen wir die zwei Zahlen a und b und ihren größten gemeinsamen Teiler d . Sei p eine beliebige Primzahl, die in den Primfaktorzerlegungen der beiden Zahlen in den Potenzen p^k bzw. p^n vorkommt, und o. B. d. A.¹² $k \leq n$. Dann ist p^k offenbar ein Teiler von a und von b und damit nach Korollar 7.5 ein Teiler von d . Andererseits kann p^{k+1} sicher *kein* Teiler von d sein, denn sonst wäre es auch ein Teiler von a . Damit ergibt sich die Primfaktorzerlegung von d : Sie besteht aus allen Primzahlen, die Teiler sowohl von a als auch von b sind, und die zugehörige Potenz ist jeweils das [Minimum](#) der entsprechenden Potenzen für a und b .

Konkretes Beispiel: 1008 ist $2^4 \cdot 3^2 \cdot 7$ und 1080 ist $2^3 \cdot 3^3 \cdot 5$. Die Primzahlen, die in beiden Zerlegungen vorkommen, sind 2 und 3. Der kleinste Exponent für 2 ist 3, der kleinste für 3 ist 2. Daher muss nach den Überlegungen aus dem letzten Absatz $2^3 \cdot 3^2 = 72$ der größte gemeinsame Teiler von 1008 und 1080 sein. Und das stimmt natürlich.

kleinstes gemeinsames
Vielfaches

$\text{kgV}(a, b)$

Aufgabe 8.7. Sie haben sicher schon einmal etwas vom *kleinsten gemeinsamen Vielfachen* zweier Zahlen a und b gehört: Zwei beliebige Zahlen a und b , die beide nicht verschwinden, haben immer mindestens ein gemeinsames positives Vielfaches, nämlich $|a| \cdot |b|$. Es gibt deshalb ein *kleinstes positives* Vielfaches, das beide gemeinsam haben. Dafür schreibt man $\text{kgV}(a, b)$. (Und ergänzend setzt man $\text{kgV}(a, b) = 0$, wenn $ab = 0$ gilt.)

Berechnen Sie das kleinste gemeinsame Vielfache für ein paar Zahlenpaare oder lassen Sie es berechnen, z. B. in [Wolfram Alpha](#) mit einer Eingabe wie `lcm(400, 225)`.¹³

Aufgabe 8.8. Erkennen Sie einen Zusammenhang zwischen $\text{ggT}(a, b)$ und $\text{kgV}(a, b)$ sowie a und b ?

Gödelisierung

★**Aufgabe 8.9.** Den [Fundamentalsatz der Arithmetik](#) kann man für eine Codierungsmethode verwenden, die *Gödelisierung* genannt wird. Der Name bezieht sich auf den

¹²Diese Abkürzung steht für *ohne Beschränkung der Allgemeinheit* und damit ist gemeint, dass man – in der Regel zur Vermeidung von Schreibarbeit – eine Einschränkung vornimmt, die aber faktisch keine ist. Im konkreten Fall gilt $k \leq n$ oder $n \leq k$. Im zweiten Fall würde man aber bis auf Vertauschung von a und b denselben Beweis erhalten und spart sich daher die Arbeit, das noch einmal aufzuschreiben.

¹³Im Englischen *least common multiple*.

österreichischen Mathematiker **Kurt Gödel**, der dieses Verfahren für den Beweis seines berühmten **Unvollständigkeitssatzes** ersann. Es ist für die Praxis nicht zu gebrauchen, für die Theorie aber sehr wichtig – auch in der theoretischen Informatik.

Man benötigt dafür einen Vorrat an Zeichen, denen positive natürliche Zahlen zugeordnet sind. Als Beispiel nehmen wir die 26 Buchstaben des lateinischen Alphabets, die von 1 (A) bis 26 (Z) nummeriert seien. Die Zeichenkette HAW würde dann als die Zahl

$$2^8 \cdot 3^1 \cdot 5^{23} = 9\,155\,273\,437\,500\,000\,000$$

codiert werden. Das Prinzip ist, dass das n -te Zeichen einer Zeichenkette auf p_n^k abgebildet wird, wobei p_n die n -te Primzahl und k die Nummer des Zeichens ist. Z. B. ist W das dritte Zeichen im Wort HAW und hat die Nummer 23 im Alphabet, also wird daraus 5^{23} . Die daraus resultierenden Potenzen werden alle multipliziert und ergeben eine Zahl.

Entscheidend ist, dass es für das Ergebnis nur eine einzige mögliche Zerlegung gibt, so dass man aus dieser die Zeichenkette eindeutig rekonstruieren kann – wenn auch ggf. mit enormem Rechenaufwand.

Codieren Sie nach diesem System das Wort **ZAPPA** und decodieren Sie die Zahl 1 464 843 750. Zusatzfrage: Ist *jede* natürliche Zahl, die größer als 1 ist, der Code eines Wortes?

★ Primzahlen im Computer

Die schon erwähnte Bibliothek SYMPY enthält diverse zahlentheoretische Funktionen, zum Beispiel diese drei:

```
from sympy import prime, isprime, factorint

prime(424242)           # n-te Primzahl
factorint(4200420042)   # Primfaktorzerlegung
isprime(4242424243)     # Test, ob Argument prim ist
```

Grundsätzlich ist allerdings zu beachten, dass das Finden von Teilern bzw. das Zerlegen in Primfaktoren ein **schwieriges mathematisches Problem** ist, für das bisher keine effizienten Algorithmen bekannt sind. Auf diesem „Mangel“ basieren diverse kryptographische Verfahren.

Auch ein verllässlicher Test, ob eine Zahl eine Primzahl ist, ist für große Argumente recht zeitaufwendig. Funktionen wie `isprime` arbeiten deshalb in der Regel **probabilistisch** und liefern Ergebnisse, die mit einer geringen Wahrscheinlichkeit auch falsch sein können. In Kapitel 9 von *Konkrete Mathematik (nicht nur) für Informatiker* wird **der in der Praxis am häufigsten eingesetzte Primzahltest** ausführlich besprochen und auch in PYTHON implementiert. Im 10. Kapitel wird das weit verbreitete **RSA-Kryptosystem** erläutert, das die gerade angesprochene Schwierigkeit bei der Faktorisierung ausnutzt. Damit Sie diese beiden Kapitel verstehen, sollten Sie aber erst den folgenden Abschnitt über modulare Arithmetik durcharbeiten.

9. Modulare Arithmetik

Die Anfang des 19. Jahrhunderts von **Carl Friedrich Gauß** entwickelte modulare Arithmetik ist in der Zahlentheorie ein häufig eingesetztes Hilfsmittel, um Rechenarbeit zu sparen. Sie kann auch beim Programmieren hilfreich sein.

Denken Sie zur Einführung kurz über die folgenden Aussagen nach, die Ihnen hoffentlich klar sind:

- Addiert man zwei gerade Zahlen, dann ist ihre Summe gerade.
- Addiert man zwei ungerade Zahlen, so ist die Summe ebenfalls gerade.
- Die Summe einer geraden und einer ungeraden Zahl ist ungerade.
- Multipliziert man zwei Zahlen, von denen mindestens eine gerade ist, dann ist das Produkt gerade.
- Multipliziert man zwei ungerade Zahlen, so ist ihr Produkt ungerade.

Die modulare Arithmetik liefert eine Begründung dafür, warum die obigen Aussagen stimmen, und verallgemeinert diese Ergebnisse.

Satz 9.1

Sei $m \in \mathbb{N}^+$. Sind a_1 und a_2 ganze Zahlen, die bei Division durch m die Reste r_1 bzw. r_2 haben, so haben $a_1 + a_2$ und $a_1 a_2$ bei Division durch m dieselben Reste wie $r_1 + r_2$ bzw. $r_1 r_2$.

Als Formel kann man das so (mit a und b statt a_1 und a_2) ausdrücken:

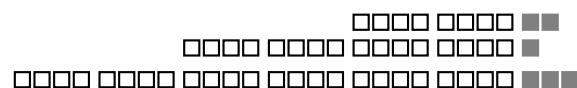
$$\begin{aligned}\forall m \in \mathbb{N}^+ \forall a, b \in \mathbb{Z} \quad (a + b) \bmod m &= ((a \bmod m) + (b \bmod m)) \bmod m \\ \forall m \in \mathbb{N}^+ \forall a, b \in \mathbb{Z} \quad (a \cdot b) \bmod m &= ((a \bmod m) \cdot (b \bmod m)) \bmod m\end{aligned}$$

Dieser Satz beinhaltet zwei grundlegende Aussagen:

- Wenn mich nur der Rest $(a + b) \bmod m$ interessiert, dann kann ich stattdessen die Reste $a \bmod m$ und $b \bmod m$ addieren und diese Summe durch m dividieren. (Ich rechne dadurch evtl. mit wesentlich kleineren Zahlen.)
- Daraus folgt, dass das Ergebnis gar nicht von a und b , sondern nur von $a \bmod m$ und $b \bmod m$ abhängt. Für eine andere Summe $a' + b'$ muss sich derselbe Rest ergeben, wenn a und a' sowie b und b' dieselben Reste haben.

Und das gilt ebenso für Produkte.

Die folgende Skizze zeigt das für die Summe von $a = 10$ und $b = 17$ für den Fall $m = 4$, wobei die Reste grau gefärbt sind.



a besteht aus zwei Blöcken zu vier Quadraten zusammen mit dem Rest von zwei Quadraten. Das ist ja die **Definition des Restes**. Analog besteht b aus vier Blöcken zu vier Quadraten und einem „Restquadrat“. Zusammen ergibt das sechs Blöcke zu vier Quadraten sowie drei „Restquadrate“. Entscheidend ist, dass man gar nicht

wissen muss, wie viele Blöcke zu vier Quadraten links stehen, wenn man nur an den drei grauen Quadraten unten rechts interessiert ist.

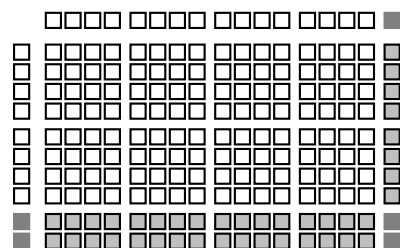
Aufgabe 9.1. Zeichnen Sie eine solche Skizze für $a = 18$ und $b = 5$. Sie sollten sehen, dass der rechte Teil (die grauen Quadrate) mit der obigen Skizze übereinstimmt.

Satz 9.1 besagt allerdings auch, dass man nicht einfach die Reste addieren darf, sondern dass man deren Summe ggf. auch noch durch m teilen muss. Das zeigt die folgende Skizze für $m = 4$, $a = 15$ und $b = 18$.



Hier ist die Summe der Reste fünf, aber der Rest der Summe ist eins, weil in der Fünf noch ein Block von vier Quadraten „steckt“.

Die Situation für das Produkt von 10 (vertikal) und 17 (horizontal) kann man ebenfalls grafisch darstellen:



Das Produkt ist 170 und es ergeben sich insgesamt $32 + 8 + 2 = 42$ Blöcke zu vier Quadraten. Beides ist aber irrelevant, wenn man nur den Rest wissen will. Den erhält man einfach durch Multiplikation der Reste eins und zwei.

Aufgabe 9.2. Zeichnen Sie so eine Skizze für $m = 4$, $a = 6$ und $b = 13$. (Die Reste sind dieselben wie eben.)

Aufgabe 9.3. Zeichnen Sie eine Skizze für $m = 5$ (!) und das Produkt von $a = 8$ und $b = 13$. Es sollte sich nun ein „Überschuss“ wie bei der Summe von 15 und 18 oben ergeben.

Modulare Arithmetik besteht nun vereinfacht gesagt darin, dass man von Anfang an nur noch mit den Resten rechnet und den „Rest“ ignoriert. Statt

$$(15 + 18) \bmod 4 = ((15 \bmod 4) + (18 \bmod 4)) \bmod 4 = (3 + 2) \bmod 4 = 1$$

werden wir in Zukunft einfach

$$[15]_4 + [18]_4 = [3]_4 + [2]_4 = [3 + 2]_4 = [5]_4 = [1]_4$$

schreiben.

Definition

Für $m \in \mathbb{N}^+$ und $a \in \mathbb{Z}$ definieren wir die **Restklasse von a modulo m** als

$$[a]_m = \{z \in \mathbb{Z} : a \bmod m = z \bmod m\}$$

und setzen die folgenden Rechenregeln für diese neuen Objekte fest:

$$\begin{aligned} [a]_m + [b]_m &= [a + b]_m \\ [a]_m \cdot [b]_m &= [a \cdot b]_m \end{aligned} \tag{9.1}$$

Dazu ein paar Anmerkungen:

- $[a]_m$ ist nach Definition die Menge aller Zahlen, die bei Division durch m denselben Rest wie a haben. Beispielsweise gehören zu $[14]_4$ unendlich viele Zahlen wie unter anderem 14, 22, 2 und -10 .
- Insbesondere ist a immer ein Element von $[a]_m$. Ferner folgt aus $b \in [a]_m$ natürlich $a \in [b]_m$ und $[a]_m = [b]_m$. Zum Beispiel gilt $[14]_4 = [-10]_4$.
- Dass $[a]_m$ rein technisch eine Menge ist, ist für das Folgende nicht so wichtig. Wesentlich ist, dass die zwischen eckigen Klammern stehenden Zahlen gewissermaßen Stellvertreter bzw. Repräsentanten für alle Zahlen mit demselben Rest bei Division durch m sind.
- Man beachte, dass die Zeichen $\cdot$ und $+$ in Gleichung (9.1) links und rechts vom Gleichheitszeichen für unterschiedliche Operationen stehen. Rechts ist die übliche **Verknüpfung** von Zahlen gemeint, während links eine Verknüpfung von Restklassen definiert wird. Dass für verschiedene Dinge dieselben Zeichen verwendet werden und sich die Bedeutung erst aus dem Kontext erschließt, ist in der Mathematik durchaus üblich – und auch **in der Informatik gang und gäbe**.
- Ausdrücke wie $[a]_m + [b]_n$ oder $[a]_m \cdot [b]_n$ sind nicht definiert und ergeben auch keinen Sinn, wenn m und n verschieden sind.
- Es ist nicht selbstverständlich, dass Gleichung (9.1) nicht zu Widersprüchen führt. Weil $[15]_4 = [-5]_4$ und $[18]_4 = [30]_4$ gilt, muss $[15]_4 + [18]_4$ dasselbe wie $[-5]_4 + [30]_4$ sein. Nach Definition kommt jedoch einmal $[33]_4$ und einmal $[25]_4$ heraus. Wir können uns nur wegen Satz 9.1 sicher sein, dass es sich in beiden Fällen um dieselbe Restklasse handelt. (Man sagt deswegen, dass die Operationen **wohldefiniert** sind.)

wohldefiniert

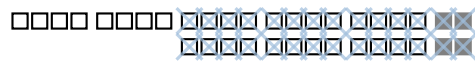
Aufgabe 9.4. Berechnen Sie zur Übung Ausdrücke wie $[25]_7 + [12]_7$ oder $[11]_6 \cdot [7]_6$. Schreiben Sie die Ergebnisse immer so, dass zwischen den eckigen Klammern ein möglichst kleiner nichtgenativer Repräsentant steht. Schreiben Sie also z. B. $[2]_9$ statt $[20]_9$ oder $[-7]_9$.

Wir halten noch einen einfachen Punkt fest, der für das Rechnen oft hilfreich ist.

Lemma 9.2

$[a]_m = [b]_m$ bzw. $a \bmod m = b \bmod m$ gilt genau dann, wenn m ein Teiler der Differenz $a - b$ ist.

Vielleicht war Ihnen das ohnehin schon klar. Ansonsten hilft evtl. die folgende Skizze für $[22]_4 = [14]_4$.



Subtrahiert man die beiden Zahlen voneinander, so heben sich die durchgestrichenen Quadrate gegenseitig auf. Insbesondere gilt das für die grauen „Restquadrate“ rechts, deren Anzahl in beiden Fällen gleich ist. Übrig bleiben können nur Blöcke zu je vier Quadraten, also eine Anzahl von Quadraten, die durch vier teilbar ist.

Die Menge $\{[a]_m : a \in \mathbb{Z}\}$ aller Restklassen modulo m wird als **Restklassenring modulo m** bezeichnet und $\mathbb{Z}/m\mathbb{Z}$ geschrieben.

Definition

Weil der Rest bei Division durch m nach Definition nur eine der Zahlen von 0 bis $m - 1$ sein kann, gibt es nur m verschiedene Restklassen modulo m . $\mathbb{Z}/m\mathbb{Z}$ ist also die Menge $\{[0]_m, [1]_m, \dots, [m-1]_m\}$. Weil diese Mengen endlich sind, kann man – anders als in $\mathbb{Z}$ oder gar $\mathbb{Q}$ – sämtliche mögliche Additionen und Multiplikationen in einer Tabelle zusammenfassen und danach das Rechnen durch das Nachschauen in der Tabelle ersetzen. Für $m = 4$ würde es so aussehen:

+	$[0]_4$	$[1]_4$	$[2]_4$	$[3]_4$
$[0]_4$	$[0]_4$	$[1]_4$	$[2]_4$	$[3]_4$
$[1]_4$	$[1]_4$	$[2]_4$	$[3]_4$	$[0]_4$
$[2]_4$	$[2]_4$	$[3]_4$	$[0]_4$	$[1]_4$
$[3]_4$	$[3]_4$	$[0]_4$	$[1]_4$	$[2]_4$

$\cdot$	$[0]_4$	$[1]_4$	$[2]_4$	$[3]_4$
$[0]_4$	$[0]_4$	$[0]_4$	$[0]_4$	$[0]_4$
$[1]_4$	$[0]_4$	$[1]_4$	$[2]_4$	$[3]_4$
$[2]_4$	$[0]_4$	$[2]_4$	$[0]_4$	$[2]_4$
$[3]_4$	$[0]_4$	$[3]_4$	$[2]_4$	$[1]_4$

Insbesondere an der Additionstabelle kann man ein sehr einfaches Schema erkennen: Jede Zeile entsteht aus der vorherigen durch zyklisches Verschieben um eine Position und das gilt auch für die Spalten. Das führt dazu, dass jedes Element in jeder Zeile bzw. Spalte genau einmal vorkommt. So eine Anordnung nennt man auch **lateinisches Quadrat**. Mit den Bezeichnungen aus dem [Abschnitt über Arithmetik](#) erkennt man:

lateinisches Quadrat

- Addition und Multiplikation sind kommutativ und assoziativ, weil sich diese Eigenschaften durch (9.1) direkt von $\mathbb{Z}$ auf $\mathbb{Z}/m\mathbb{Z}$ übertragen. Zum Beispiel:

$$[a]_m + [b]_m = [a + b]_m = [b + a]_m = [b]_m + [a]_m \quad (9.2)$$

- $[0]_m$ ist beim Addieren in $\mathbb{Z}/m\mathbb{Z}$ das neutrale Element, d. h., es gilt immer $[a]_m + [0]_m = [a]_m$.
- Jedes Element $[a]_m$ von $\mathbb{Z}/m\mathbb{Z}$ hat ein inverses Element bzgl. der Addition, nämlich $[-a]_m = [m - a]_m$. Invers zu $[7]_{10}$ ist z. B. $[3]_{10}$.

Für diese inversen Elemente schreiben wir wie üblich $-[a]_m$ und verwenden wie bei den ganzen Zahlen $[a]_m - [b]_m$ als Abkürzung für $[a]_m + (-[b]_m)$. Praktischerweise gilt dann

$$[a]_m - [b]_m = [a - b]_m. \quad (9.3)$$

Aufgabe 9.5. Können Sie die einzelnen Schritte in (9.2) und die jeweilige Bedeutung erklären?

Aufgabe 9.6. Üben Sie wie in Aufgabe 9.4 das Subtrahieren von Restklassen und das Berechnen von inversen Elementen.

Rechnen in $\mathbb{Z}/n\mathbb{Z}$

Für die weitere Arbeit mit Restklassen führen wir nun eine Konvention ein, um Schreibarbeit zu sparen: Wenn durch eine Formulierung wie „wir rechnen in $\mathbb{Z}/n\mathbb{Z}$ “ oder „wir rechnen modulo n “ klar ist, dass mit einem festen n gerechnet wird, schreiben wir statt $[a]_n$ abkürzend die Zahl $a \bmod n$. $\mathbb{Z}/n\mathbb{Z}$ wird dadurch zu $\{0, \dots, n-1\}$.

Beim Rechnen modulo 5 haben wir es dann also z. B. mit der Menge $\{0, 1, 2, 3, 4\}$ zu tun. Statt $[4]_5$ schreiben wir nun also einfach 4. Beachten Sie, dass damit *nicht* die ganze Zahl 4 gemeint ist! Insbesondere gelten dann ganz neue und „seltsame“ Rechenregeln: $4 + 3$ ist nun 2, denn es gilt ja $[4]_5 + [3]_5 = [2]_5$. Und $4 + 3$ ist *nicht* 7, obwohl auch $[4]_5 + [3]_5 = [7]_5$ nicht falsch ist. Nach der obigen Konvention wird für $[7]_5$ jedoch 2 geschrieben.

Aufgabe 9.7. Probieren wir es gleich einmal aus. Berechnen Sie in $\mathbb{Z}/8\mathbb{Z}$ die Werte $7 + 6$, $4 - 6$ und $4 \cdot 6$.

Aufgabe 9.8. Was haben die Aussagen des einführenden Beispiels von Seite 68 über gerade und ungerade Zahlen mit modularer Arithmetik zu tun und wie kann man ihre Korrektheit begründen?

★Aufgabe 9.9. An einer Tafel stehen die natürlichen Zahlen von 1 bis einschließlich 2025. Man darf irgend zwei Zahlen wegwischen und dafür ihre Differenz anschreiben. Wiederholt man diesen Vorgang genügend oft, so bleibt an der Tafel schließlich nur noch eine Zahl stehen. Diese Zahl ist garantiert ungerade. Können Sie begründen, warum das so ist?

Aufgabe 9.10. Beim Programmieren haben Sie es auch manchmal (häufig eher ungewollt) mit modularer Arithmetik zu tun. Wenn ich auf meinem Rechner mit der folgenden C-Funktion die Fakultät von 13 berechnen will, dann erhalte ich als Ausgabe die Zahl 1932053504.

```
int fact (int n) {
    int f = 1;
    while (n > 0) {
        f = f * n;
        n = n - 1;
    }
    return f;
}
```

Aber die korrekte Antwort wäre 6227020800 gewesen. Woran liegt das?

Aufgabe 9.11. Nebenbei bemerkt kann man übrigens *ohne Rechnen* sofort sehen, dass die Antwort 1932053504 aus Aufgabe 9.10 nicht richtig sein kann. Wieso?

Aufgabe 9.12. Wenn man modulo 5 rechnet, kann man 7 durch 2 ersetzen. Daher müssten 2^2 und 2^7 denselben Rest bei Division durch 5 haben, oder? Das stimmt aber nicht! (Rechnen Sie nach.) Was ist an dieser Argumentation falsch?

Aufgabe 9.13. Erinnern Sie sich an die Teilbarkeitsregeln mit den **Quersummen** von Seite 57. Begründen Sie an einem konkreten Beispiel wie 7164, warum das klappt, indem Sie in $\mathbb{Z}_9$ bzw. $\mathbb{Z}_3$ rechnen. (Welche Reste erhalten Sie, wenn Sie 1, 10, 100, 1000 und so weiter durch 9 bzw. 3 teilen?)

Aufgabe 9.14. Um auf Teilbarkeit durch 11 zu testen, geht man wie bei der Teilbarkeit durch 9 vor, berechnet jedoch die *alternierende Quersumme*, indem man – bei der kleinsten Stelle beginnend – die Ziffern abwechselnd addiert und subtrahiert. Die alternierende Quersumme von 97364 ist also $4 - 6 + 3 - 7 + 9 = 3$. Und weil 3 nicht durch 11 teilbar ist, ist auch 97364 nicht durch 11 teilbar. Finden Sie analog zu Aufgabe 9.13 eine Begründung dafür, warum diese Regel funktioniert.

alternierende
Quersumme

★Aufgabe 9.15. Für die Sieben gibt es eine weitere Teilbarkeitsregel, die allerdings etwas komplizierter als die bisherigen ist. Um festzustellen, ob eine Zahl durch 7 teilbar ist, verdoppeln Sie die letzte Ziffer und subtrahieren das Ergebnis vom Rest der Zahl (ohne die letzte Ziffer). Wiederholen Sie das so lange, bis Sie dem Ergebnis ansehen können, ob es durch 7 teilbar ist. Das ist genau dann der Fall, wenn auch die ursprüngliche Zahl durch 7 teilbar ist.

Um z. B. zu sehen, ob 35581 durch 7 teilbar ist, berechnen wir $3558 - 2 \cdot 1 = 3556$. Aus diesem Ergebnis machen wir $355 - 2 \cdot 6 = 343$. Daraus wird $34 - 2 \cdot 3 = 28$. Und dass 28 durch 7 teilbar ist, wissen wir. Darum ist auch 35581 durch 7 teilbar.

Können Sie begründen, warum das immer klappt?

Bisher haben wir uns noch gar nicht mit der Division beschäftigt. Klar ist jedenfalls, dass wir nicht analog zu (9.1) und (9.3) $[a]_m / [b]_m = [a/b]_m$ schreiben können.

Aufgabe 9.16. Warum ist das klar?

Sinnvoll wäre lediglich (analog zur Herleitung der Subtraktion weiter oben) eine Definition der Division als Abkürzung für die Multiplikation mit dem Kehrwert. (Siehe Seite 42.) Dass das nicht immer klappen kann, zeigt bereits die Multiplikationstabelle von $\mathbb{Z}/6\mathbb{Z}$:

·	0	1	2	3	4	5
0	0	0	0	0	0	0
1	0	1	2	3	4	5
2	0	2	4	0	2	4
3	0	3	0	3	0	3
4	0	4	2	0	4	2
5	0	5	4	3	2	1

Die ausgegrauten Teile können wir zwar ignorieren, weil Multiplikation mit null auch in der modularen Arithmetik immer das Produkt null liefert. Aber in der Resttabelle müsste in jeder Zeile und Spalte eine Eins stehen,¹⁴ damit jedes Element von $\mathbb{Z}/6\mathbb{Z}$ außer null einen Kehrwert hätte. Das ist jedoch nicht der Fall.

¹⁴Die Eins ist offenbar das neutrale Element der Multiplikation in jedem Restklassenring.

Aufgabe 9.17. Welche Elemente von $\mathbb{Z}/6\mathbb{Z}$ haben einen Kehrwert?

Aufgabe 9.18. Stellen Sie Multiplikationstabellen für $\mathbb{Z}/5\mathbb{Z}$ und $\mathbb{Z}/7\mathbb{Z}$ auf. Können Sie einen signifikanten Unterschied zum Fall $\mathbb{Z}/6\mathbb{Z}$ erkennen?

Man kann der Tabelle für $\mathbb{Z}/6\mathbb{Z}$ noch weitere „Merkwürdigkeiten“ entnehmen, die wir von der **üblichen Arithmetik** nicht gewohnt sind:

- Man kann die Gleichung $x \cdot 2 = 3$ nicht lösen.
- Es gilt $5 \cdot 2 = 2 \cdot 2$, obwohl 5 und 2 zwei verschiedene Elemente von $\mathbb{Z}/6\mathbb{Z}$ sind.
- Es gilt $3 \cdot 2 = 0$, obwohl beide Faktoren von null verschieden sind.¹⁵

Alle drei Beispiele lassen sich darauf zurückführen (siehe Aufgabe 9.22), dass 2 in $\mathbb{Z}/6\mathbb{Z}$ keinen Kehrwert hat. Wenn Sie Aufgabe 9.18 bearbeitet haben, werden Sie jedoch gesehen haben, dass so etwas in $\mathbb{Z}/5\mathbb{Z}$ und $\mathbb{Z}/7\mathbb{Z}$ nicht passiert.

Satz 9.3

Im Restklassenring $\mathbb{Z}/n\mathbb{Z}$ hat genau dann jedes Element außer null einen Kehrwert, wenn n eine Primzahl ist.

Restklassenkörper
 $\mathbb{Z}_p$

Ist p eine Primzahl, dann schreibt man auch $\mathbb{Z}_p$ statt $\mathbb{Z}/p\mathbb{Z}$ und spricht vom *Restklassenkörper modulo p* . Wir würden also beispielsweise $\mathbb{Z}_5$ oder $\mathbb{Z}_{19}$ schreiben, aber *nicht* $\mathbb{Z}_{16}$.

Wieso klappt das mit Primzahlen und sonst nicht? Schauen wir uns das konkret für die Primzahl $n = 97$ an und versuchen wir, einen Kehrwert für das Element 67 im Restklassenkörper $\mathbb{Z}_{97}$ zu finden. Die zu lösende Gleichung ist in ausführlicher Form $[67]_{97} \cdot [x]_{97} = [1]_{97}$. „Übersetzt“ man das zurück in die ganzen Zahlen, so suchen wir eine ganze Zahl $x \in \mathbb{Z}$ mit $67x \bmod 97 = 1$. Das bedeutet, dass es ein $k \in \mathbb{Z}$ mit $67x = 97k + 1$ geben muss.

Das **Lemma von Bézout** besagt, dass es solche Zahlen x und k tatsächlich gibt, weil $a = 67$ und 97 teilerfremd sind. Und wir hatten in diesem Zusammenhang gelernt, dass man sie mithilfe des **euklidischen Algorithmus** findet:

A	B
$97 = n$	$67 = a$
$30 = n - a$	$67 = a$
$30 = n - a$	$7 = a - 2(n - a) = 3a - 2n$
$2 = (n - a) - 4(3a - 2n) = 9n - 13a$	$7 = 3a - 2n$
$2 = 9n - 13a$	$1 = (3a - 2n) - 3(9n - 13a) = 42a - 29n$

Der Prozess wurde in diesem Fall etwas abgekürzt. (Siehe dazu Aufgabe 7.14.) Entscheidend ist die Zahl 42, die am Ende als Faktor vor a steht (während der zweite Faktor 29 nicht weiter interessant ist): $1 = 42a - 29n = 42 \cdot 67 - 29 \cdot 97$ bzw. $42 \cdot 67 = 29 \cdot 97 + 1$ bedeutet $42 \cdot 67 \bmod 97 = 1$ oder $[42]_{97} \cdot [67]_{97} = [1]_{97}$. 42 ist die gesuchte Zahl x . (Rechnen Sie nach, dass in $\mathbb{Z}_{97}$ tatsächlich $42 \cdot 67 = 1$ gilt.)

¹⁵Siehe Korollar 5.2.

Der entscheidende Punkt war dabei, dass 67 und 97 teilerfremd sind. Damit das *immer* klappt, müssen alle positiven natürlichen Zahlen unterhalb von n teilerfremd zu n sein. Und das ist genau dann der Fall, wenn n eine Primzahl ist.

Haben wir es mit einem Restklassenkörper zu tun, dann können wir also Division wie auf Seite 42 als „Shortcut“ definieren. In $\mathbb{Z}_{97}$ würden wir für den Kehrwert von 67 beispielsweise $1/67 = 42$ schreiben und könnten Berechnungen wie

$$13/67 = 13 \cdot (1/67) = 13 \cdot 42 = 61$$

durchführen.

Aufgabe 9.19. Berechnen Sie so wie eben $1/106$ in $\mathbb{Z}_{173}$. Ermitteln Sie außerdem den Kehrwert von 5 in $\mathbb{Z}_{13}$.

Aufgabe 9.20. Stellen Sie sich selbst entsprechende Aufgaben. Sie können Ihre Ergebnisse überprüfen, indem Sie in [Wolfram Alpha](#) so etwas wie

inverse of 67 modulo 97

eingeben.

Aufgabe 9.21. Was ist $4/5$ in $\mathbb{Z}_{13}$?

★**Aufgabe 9.22.** Können Sie begründen, warum in einem Restklassenkörper keine der „Merkwürdigkeiten“ von Seite 74 eintreten kann?

Aufgabe 9.23. Berechnen Sie $(-1) \cdot (-1)$ in $\mathbb{Z}_4$, in $\mathbb{Z}_7$ und in $\mathbb{Z}_{13}$. Können Sie ein Muster erkennen?

Aufgabe 9.24. Kann man natürliche Zahlen a und b finden, so dass die Gleichung

$$a^2 + b^2 = 10000003 \tag{9.4}$$

gilt? Die Antwort ist nein. Man könnte ein Programm schreiben, um das durch systematisches Probieren herauszufinden (machen Sie das ruhig einmal!), aber man kann die Frage auch fast ohne Rechnen beantworten, indem man modulo 4 rechnet. Nämlich wie?

Aufgabe 9.25. X und Y spielen ein Spiel, bei dem sie abwechselnd eine Dezimalziffer (also eine der Ziffern von 0 bis inklusive 9) in eines der folgenden Felder eintragen, sofern es noch frei ist. X beginnt.

A	B	C	D	E	F

Nach sechs Zügen ist das Spiel beendet und die sechs Ziffern ergeben eine sechstellige Zahl m zwischen 000000 und 999999. Vor Beginn des Spiels wurde zufällig eine Zahl $n \in [10] \setminus \{7\}$ ausgewählt. Y hat gewonnen, wenn m am Ende durch n teilbar ist, ansonsten hat X gewonnen. Für welche n hat X eine sichere Gewinnstrategie, für welche Y? Und wie sehen diese Strategien aus?

★**Aufgabe 9.26.** In Aufgabe 9.25 wurde der Fall $n = 7$ ausgelassen, weil er ein kleines bisschen komplizierter ist. Können Sie auch für diese Zahl eine Gewinnstrategie für X oder Y angeben?

Zum Abschluss dieses Abschnitts noch einen Satz, der in der Zahlentheorie häufig verwendet wird. Er besagt, dass in jedem Restklassenkörper für einen bestimmten Exponenten die Potenz immer dieselbe ist.

Satz 9.4

Kleiner Satz von Fermat

Ist p prim, so gilt in $\mathbb{Z}_p$ die Gleichung $a^{p-1} = 1$ für alle $a \in \mathbb{Z}_p \setminus \{0\}$.

Dieser Satz ist die Grundlage für diverse in der Praxis relevante **Primzahltests**, von denen einige in der Kryptographie-Vorlesung genauer betrachtet werden. (Siehe auch Seite 67.) Seinen etwas seltsamen Namen trägt er übrigens, weil es auch einen „großen“ Satz von **Fermat** gibt. Mehr dazu in meinem Video *Der Große Satz von Fermat*.

Dass der „kleine“ Satz von Fermat wahr ist, kann man sich anhand eines Beispiels gut klarmachen. Nehmen wir etwa $p = 7$ und $a = 3$. Nach Aufgabe 9.22 sind die Produkte $1 \cdot 3$, $2 \cdot 3$ und so weiter bis $6 \cdot 3$ sechs unterschiedliche Elemente von $\mathbb{Z}_7 \setminus \{0\}$. Es muss sich somit um die Werte 1 bis 6 handeln. Da die Multiplikation in Restklassenringen kommutativ ist, impliziert das

$$(1 \cdot 3) \cdot (2 \cdot 3) \cdot (3 \cdot 3) \cdot (4 \cdot 3) \cdot (5 \cdot 3) \cdot (6 \cdot 3) = 1 \cdot 2 \cdot 3 \cdot 4 \cdot 5 \cdot 6.$$

Dividiert man nun beide Seiten dieser Gleichung nacheinander durch 2, 3 und so weiter bis 6, so steht links 3^6 und rechts nur noch die Eins. Das war's schon.

Aufgabe 9.27. Wenn in einem Restklassenring $\mathbb{Z}/n\mathbb{Z}$ für ein Element $a \in (\mathbb{Z}/n\mathbb{Z}) \setminus \{0\}$ die Gleichung $a^{n-1} = 1$ gilt, dann folgt daraus *nicht* notwendig, dass n prim ist. Finden Sie ein nichttriviales Beispiel dafür.

Aufgabe 9.28. Wir wissen bereits, dass man Kehrwerte in Restklassenkörpern mittels des erweiterten euklidischen Algorithmus berechnen kann. Der kleine Satz von Fermat liefert noch eine weitere (zumindest theoretische) Möglichkeit zur Berechnung von Kehrwerten. Sehen Sie sie?

Potenzen in Restklassenringen (siehe dazu Aufgabe 9.12) lassen sich oft auch dann erstaunlich einfach berechnen, wenn der Exponent sehr groß ist. Wollen wir beispielsweise 10^{200} in $\mathbb{Z}_{11}$ berechnen, so können wir ausnutzen, dass $10 = -1$ gilt. Nach Aufgabe 9.23 gilt $(-1)^2 = 1$ und damit haben wir

$$10^{200} = (10^2)^{100} = ((-1)^2)^{100} = 1^{100} = 1.$$

Aufgabe 9.29. Berechnen Sie mit möglichst geringem Rechenaufwand 4^{100} in $\mathbb{Z}_{13}$. (Hinweis: Berechnen Sie zuerst 4^3 .)

Da Satz 9.4 wichtig für die Kryptographie ist, braucht man ein Verfahren, mit dem Computer Potenzen der Form a^b auch ohne „Tricks“ schnell ausrechnen können, wenn a und b sehr große (natürliche) Zahlen sind.¹⁶ Das ist die sogenannte *binäre Exponentiation* (die im Englischen *Square and Multiply* genannt wird). Dabei nutzt man (v) von Lemma 5.8 aus, und zwar speziell $a^{2n} = (a^n)^2$. Während man für den linken Term $2n - 1$ Multiplikationen durchführen muss, sind es rechts nur n . Das funktioniert zwar nur für gerade Exponenten, aber für ungerade kann man (nach demselben Lemma) einfach a ausklammern: $a^{2n+1} = a \cdot a^{2n}$ und käme auf $n + 1$ statt $2n$ Multiplikationen. Die tatsächliche Ersparnis ist sogar noch viel größer, weil man dasselbe Verfahren *rekursiv* auch auf a^n anwenden kann. Beispielsweise ist a^{42} dasselbe wie $(a^{21})^2$. a^{21} kann man als $a \cdot a^{20}$ schreiben und dann statt a^{20} besser $(a^{10})^2$ berechnen. a^{10} ersetzt man durch $(a^5)^2$ und so weiter.

binäre Exponentiation

Aufgabe 9.30. In der modularen Arithmetik kann man zudem ausnutzen, dass auch bei hohen Exponenten alle Zwischenergebnisse „klein“ sind. Berechnen Sie 5^{37} in $\mathbb{Z}/9\mathbb{Z}$ mithilfe der Technik der binären Exponentiation, ohne einen Taschenrechner oder andere elektronische Rechenhilfen zu verwenden.

★Aufgabe 9.31. Können Sie eine von b abhängende Obergrenze für die Anzahl der Multiplikation angeben, die man benötigt, um a^b mit binärer Exponentiation auszurechnen? [Tipp: Schreiben Sie b *binär* auf.]

Siehe dazu eine Beispielimplementierung im folgenden *Abschnitt über modulare Arithmetik im Computer*. Weitere Informationen und Beispiele im Video *Binäre Exponentiation*.

Restklassenkörper werden in vielen *Fehlererkennungs- und Fehlerkorrekturverfahren* eingesetzt. So wird die Prüfziffer für Buchnummern im System *ISBN-10* z. B. in $\mathbb{Z}_{11}$ berechnet.¹⁷ Und das Verfahren, das verhindern soll, dass Sie beim Online-Banking aus Versehen eine falsche *IBAN* eingeben, rechnet in $\mathbb{Z}_{97}$.

Dass $\mathbb{Z}/n\mathbb{Z}$ kein *Körper* ist, wenn n keine Primzahl ist, heißt übrigens nicht notwendig, dass es keinen Körper mit genau n Elementen gibt. Es gibt beispielsweise Körper mit 9, 32 oder 125 Elementen. Solche sogenannten *Galoiskörper* werden unter anderem für die automatische Korrektur von Fehlern in digitalen Medien eingesetzt. Mehr dazu im Video *Fehlerkorrektur mit Reed-Solomon-Codes*.

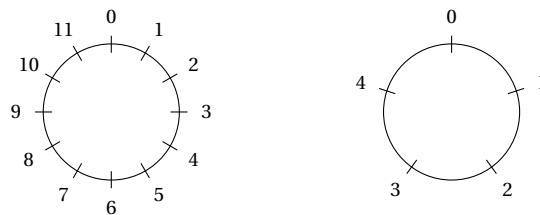
Bemerkungen

Es sei noch einmal daran erinnert, dass nach der Konvention von Seite 72 Rechnen „in $\mathbb{Z}/n\mathbb{Z}$ “ (bzw. „in $\mathbb{Z}_n$ “, wenn n prim ist) bedeutet, dass es nur noch die „Zahlen“ 0 bis $n - 1$ gibt. Wenn in einer Klausuraufgabe z. B. steht, dass in $\mathbb{Z}_7$ gerechnet wird, und Sie als Ergebnis Terme wie -3 , 11 oder $6/5$ aufschreiben, so ist das keine akzeptable Antwort. Es gilt zwar $[-3]_7 = [11]_7 = [6]_7/[5]_7 = [4]_7$, aber richtig wäre in diesem Fall nur die Antwort 4 gewesen.

¹⁶Die Berechnung von Potenzen läuft intern übrigens völlig anders ab, wenn mit *Fließkommazahlen* gearbeitet wird. Dazu kommen wir später.

¹⁷Siehe dazu Kapitel 6 in *Konkrete Mathematik (nicht nur) für Informatiker*.

Modulare Arithmetik ist gewöhnungsbedürftig, aber Sie praktizieren sie eigentlich schon seit Jahren. Es ist die übliche Art und Weise, in der wir alle auf dem Ziffernblatt einer Uhr rechnen. Wenn jemand Ihnen sagt, dass er um zehn Uhr morgens ankam und fünf Stunden blieb, dann wissen Sie ohne großes Nachdenken, dass er um drei Uhr nachmittags gegangen ist. Implizit haben Sie $10 + 5 = 15$ gerechnet und dann 12 abgezogen. Oder Sie haben es „visuell“ ausgerechnet, indem Sie den „Zeiger“ auf 10 stellten und dann fünf Schritte vorwärts bewegten. Sie rechnen dabei *modulo 12*. Wenn Sie stattdessen modulo 5 rechnen, dann ist das dasselbe Prinzip, nur mit einer anderen „Uhr“.



Wie schon mehrfach erwähnt ist die Zahlentheorie wichtig für die Kryptographie. Dort wird in der Praxis mit viel größeren Zahlen als in den Beispielen in diesem Kapitel gerechnet. Diese sind teilweise so groß, dass man allein für das Hinschreiben der Zahl eine ganze Buchseite bräuchte. Bei Zahlen dieser Größenordnung kann der Wechsel von einem langsamen und zu einem schnellen Algorithmus bedeuten, dass man wenige Sekunden statt vieler Jahre braucht. (Siehe z. B. die Lösung von Aufgabe 9.31.)

★ Modulare Arithmetik im Computer

Viele Berechnungen, die mit modularer Arithmetik zu tun haben, lassen sich in PYTHON einfach mit dem [schon erwähnten](#) Operator `%` durchführen. Den Test, ob zwei Zahlen a und b bei Division durch $m \in \mathbb{N}^+$ denselben Rest haben, kann man nach Lemma 9.2 auf zwei Arten implementieren:

```
sameClass = lambda a,b,m: a%m==b%m           # Variante 1
sameClass = lambda a,b,m: divides(m,a-b)      # Variante 2
```

Dabei wurde für die zweite Variante die [bereits definierte](#) Funktion `divides` verwendet.

SYMPY hat eine spezielle Funktion zur Berechnung der Kehrwerte:

```
from sympy import mod_inverse
mod_inverse(67,97)
```

Außerdem hat PYTHON eine Standardfunktion `pow` für Potenzen (nicht zu verwechseln mit der [gleichnamigen Funktion aus der MATH-Bibliothek](#)), die man mit einem dritten Argument dazu bringen kann, modular zu rechnen:

```
4**100, pow(4,100), pow(4,100,13)
```

Durch den Einsatz der binären Exponentiation ist diese Funktion auch für große Exponenten sehr schnell. Man könnte das aber auch selbst mit geringem Aufwand realisieren:

```
def Pow(a,b,n):
    if b==0: return 1
    if b%2==1: return a*Pow(a,b-1,n) % n    # multiply
    return Pow(a,b//2,n)**2 % n            # square
```

10. Der chinesische Restsatz

In diesem Abschnitt behandeln wir kurz einen Algorithmus, der in der Kryptographie an verschiedenen Stellen benutzt wird.¹⁸ Das Verfahren basiert auf Ideen, die schon dem chinesischen Mathematiker [Sun Zi](#) im dritten Jahrhundert bekannt waren.

$a \equiv b \pmod{m}$ ist lediglich eine andere, in der Zahlentheorie häufig verwendete, Schreibweise für $[a]_m = [b]_m$ bzw. $a \bmod m = b \bmod m$.¹⁹ Solche Aussagen werden auch als **Kongruenzen** bezeichnet.

Definition

Chinesischer Restsatz

Ein Kongruenzsystem der Form

$$\begin{aligned} x &\equiv a_1 \pmod{m_1} \\ x &\equiv a_2 \pmod{m_2} \\ &\vdots \\ x &\equiv a_n \pmod{m_n} \end{aligned}$$

hat für beliebige vorgegebene natürliche Zahlen a_i eine Lösung, wenn m_1 bis m_n paarweise²⁰ teilerfremde natürlichen Zahlen sind.

Satz 10.1

Bei einem *Kongruenzsystem* wird also – ähnlich wie bei einem (linearen) Gleichungssystem (das Sie aus der Schule kennen) – eine Zahl x gesucht, die alle Kongruenzen gleichzeitig simultan erfüllt. Und natürlich geht es wie im gesamten

Kongruenzsystem

¹⁸Unter anderem für [Geheimnisteilungsverfahren](#) sowie für die Steigerung der Effizienz des [RSA-Kryptosystems](#), das in Kapitel 10 von [Konkrete Mathematik \(nicht nur\) für Informatiker](#) ausführlich beschrieben wird.

¹⁹Man muss allerdings beachten, dass $\bmod$ in einem Ausdruck wie $a \bmod b$ ein zweistelliger [Operator](#) wie $+$ ist, während die Zeichenfolge $\pmod{m}$ nur *zusammen* mit dem Ausdruck $a \equiv b$ davor Sinn ergibt.

²⁰Das Adverb *paarweise* ist typischer „Mathematikerschnack“. An dieser Stelle ist es so gemeint, dass immer dann, wenn man sich zwei Zahlen (also ein *Paar*) herausgreift, sie teilerfremd sind.

paarweise

Kapitel um eine ganze Zahl. (Sonst würden die Kongruenzen auch gar keinen Sinn ergeben.)

Die Lösungsidee kann man sich an dem einfachen System

$$\begin{aligned} x &\equiv 1 \pmod{3} \\ x &\equiv 2 \pmod{5} \end{aligned} \tag{10.1}$$

klarmachen. Markieren Sie ein paar Lösungen der ersten Kongruenz auf der Zahlengeraden: 1, 4, 7, 10, −2, −5 und so weiter. Markieren Sie dann ein paar Lösungen der zweiten Kongruenz in einer anderen Farbe. Sie werden schnell merken, dass es Zahlen gibt, die beide Kongruenzen lösen, zum Beispiel 7. Aber wie kommt man ohne Probieren auf so eine Lösung?

Der Trick ist, nach einer Lösung x_1 der ersten Kongruenz zu suchen, die durch 5 teilbar ist. Warum? Weil wir die zu einer Lösung der zweiten Kongruenz addieren könnten und das Ergebnis wieder eine Lösung der zweiten Kongruenz wäre. Ebenso können wir eine Lösung x_2 der zweiten Kongruenz, die durch 3 teilbar ist, „gefährlos“ zu einer Lösung der ersten Kongruenz addieren. Insbesondere muss, wenn wir solche Zahlen finden, $x_1 + x_2$ eine Lösung des Systems (10.1) sein. Und das war's schon! (Im konkreten Fall könnten $x_1 = 10$ und $x_2 = 12$ diese Rollen spielen und die Lösung 22 liefern.)

Nun ein Beispiel mit drei Zahlen:

$$\begin{aligned} x &\equiv 2 \pmod{5} \\ x &\equiv 1 \pmod{6} \\ x &\equiv 3 \pmod{11} \end{aligned}$$

Wir überprüfen zunächst die paarweise Teilerfremdheit: 5 und 6 sind teilerfremd, 5 und 11 sind teilerfremd und 6 und 11 ebenfalls. Der chinesische Restsatz garantiert also die Existenz einer Lösung. An der folgenden Vorgehensweise kann man erkennen, warum das so ist und dass es auch mit mehr als drei Kongruenzen funktionieren würde.

Weil die drei Zahlen paarweise teilerfremd sind, ist jede von ihnen auch teilerfremd zum Produkt der anderen. Wären beispielsweise 5 und das Produkt $6 \cdot 11 = 66$ nicht teilerfremd, so gäbe es nach dem [Lemma von Bézout](#) Zahlen k_1 und k_2 mit

$$k_1 \cdot 5 + k_2 \cdot 66 > 1.$$

Diese Gleichung würde aber gleichzeitig nach [Lemma 7.3](#) ein Beleg dafür sein, dass 5 und 6 nicht teilerfremd sind, weil man den rechten Summanden auch als $(k_2 \cdot 11) \cdot 6$ lesen kann. Und das wäre ein Widerspruch.

Wiederum nach dem [Lemma von Bézout](#) gibt es also Zahlen r_1 und s_1 mit

$$r_1 \cdot 5 + s_1 \cdot 66 = 1.$$

Konkret könnten das $r_1 = -13$ und $s_1 = 1$ sein. Multipliziert man diese Gleichung mit der Zahl 2 aus der Kongruenz $x \equiv 2 \pmod{5}$, dann erhält man

$$-2 \cdot 13 \cdot 5 + \underbrace{2 \cdot 1 \cdot 66}_{x_1} = 2. \quad (10.2)$$

Die hervorgehobene Zahl $x_1 = 132$ ist nach Konstruktion einerseits ein Vielfaches von 6 und von 11 und andererseits löst sie die erste Kongruenz.

Ebenso findet man eine Zahl x_2 , die Vielfaches von 5 und 11 ist und gleichzeitig die zweite Kongruenz löst, sowie ein Vielfaches x_3 von 5 und 6, das die dritte Kongruenz löst. Das könnten z. B. $x_2 = 55$ und $x_3 = 300$ sein.

Die Summe $x = x_1 + x_2 + x_3 = 132 + 55 + 300 = 487$ löst offenbar das Kongruenzsystem: x_1 löst die erste Kongruenz und die anderen beiden Summanden sind Vielfache von 5, also löst x die erste Kongruenz. Analog überzeugt man sich, dass x die anderen beiden Kongruenzen löst.

Nachdem man sich einmal überzeugt hat, dass es prinzipiell klappt, kann man für vergleichsweise kleine Zahlen das Berechnen von x_1 , x_2 und so weiter übrigens durch Probieren ersetzen. Gleichung (10.2) besagt beispielsweise, dass wir eine Zahl k_1 mit $[66k_1]_5 = [2]_5$ suchen. Weil wir modulo 5 rechnen, vereinfacht sich das zu $[1k_1]_5 = [2]_5$. Man sieht hier sofort, dass $k_1 = 2$ gelten muss. Aber selbst dann, wenn man das nicht sehen würde, kämen für k_1 wegen der modularen Arithmetik nur vier Werte infrage – und genau einer von denen muss es sein.

Aufgabe 10.1. Lösen Sie das folgende Kongruenzsystem, indem Sie den Schritten des obigen Beispiels folgen:

$$x \equiv 1 \pmod{3}$$

$$x \equiv 2 \pmod{5}$$

$$x \equiv 4 \pmod{7}$$

Aufgabe 10.2. Stellen Sie sich selbst solche Aufgaben und lösen Sie sie – auch welche mit mehr als drei Kongruenzen. Man kann bei solchen Aufgaben offensichtlich sofort eine Probe machen und damit die eigenen Ergebnisse überprüfen. Wenn Sie Ihren Rechenkünsten nicht trauen, können Sie aber auch in [Wolfram Alpha](#) so etwas wie `ChineseRemainder[{1,2,4},{3,5,7}]` eingeben.

Aufgabe 10.3. Ihnen wird bereits aufgefallen sein, dass Kongruenzsysteme, die die Voraussetzungen des chinesischen Restsatzes erfüllen, mehr als eine Lösung haben. Wie viele sind es und welcher Zusammenhang besteht zwischen verschiedenen Lösungen?

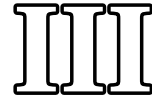
Aufgabe 10.4. Auf die Voraussetzung im chinesischen Restsatz, dass die m_i paarweise teilerfremd sind, kann man nicht verzichten. Begründen Sie das, indem Sie ein Kongruenzsystem angeben, für das man garantiert keine Lösung finden kann.

★ Der chinesische Restsatz im Computer

In der schon mehrfach erwähnten Bibliothek SYMPY gibt es eine Funktion für den chinesischen Restsatz:

```
from sympy.ntheory.modular import crt  
crt([3,5,7],[1,2,4])
```

Sie können sich sicher denken, welche Bedeutung der zweite Rückgabewert hat.



Erweiterung der Zahlenwelt

Bisher sind wir nur bis zu den **rationalen Zahlen** gekommen. Für wesentliche Teile der Mathematik reichen diese jedoch nicht aus. In diesem Kapitel führen wir zwei weitere wichtige Klassen von Zahlen ein, ohne die viele Anwendungen in den Naturwissenschaften und der Technik undenkbar wären.

11. Ungleichungen und die Zahlengerade

Wir haben Begriffe wie „positiv“ oder „kleiner“ bereits informell verwendet. Auf den folgenden Seiten nehmen wir diese Konzepte genauer unter die Lupe. Wir fangen damit an, sie präzise zu definieren. Obwohl Sie selbstverständlich wissen sollten, was gemeint ist, können diese Definition gegebenenfalls bei der Klärung von Zweifelsfällen hilfreich sein. Außerdem sind die Konzepte dieses Abschnitts relevant für das Verständnis der Rolle der **reellen Zahlen**.

Rationale Zahlen, die man in der Form m/n mit $m, n \in \mathbb{N}^+$ schreiben kann, bezeichnen wir als **positiv**. Zahlen der Form $-q$ werden **negativ** genannt, wenn q positiv ist. Für die Menge der positiven rationalen Zahlen schreiben wir $\mathbb{Q}^+$, für die Menge der negativen $\mathbb{Q}^-$.

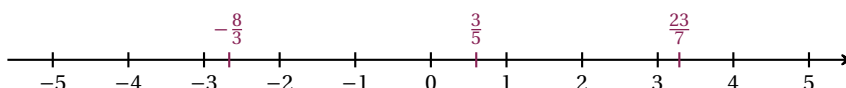
Definition

Damit zerfällt die Menge der rationalen Zahlen folgendermaßen in drei paarweise *disjunkte* Mengen:¹ $\mathbb{Q} = \mathbb{Q}^+ \dot{\cup} \mathbb{Q}^- \dot{\cup} \{0\}$.

disjunkt

Für $a, b \in \mathbb{Q}$ ist $a < b$ („ a ist kleiner als b “) eine Abkürzung für $b - a \in \mathbb{Q}^+$.
 $a > b$ ist eine andere Schreibweise für $b < a$.

Definition



¹Zwei Mengen A und B werden als *disjunkt* bezeichnet, wenn $A \cap B = \emptyset$ gilt. Das Symbol $\dot{\cup}$ steht für die **Vereinigung** disjunkter Mengen.

Eine hilfreiche Vorstellung ist die **bereits erwähnte** und aus der Schule bekannte **Zahlengerade**, von der oben ein kleiner Ausschnitt gezeigt wird:

- Jede rationale Zahl entspricht einem Punkt auf dieser Geraden.
- Die positiven Zahlen liegen rechts von der Null, die negativen links von ihr.
- Jede Zahl a liegt so, dass ihr Abstand von der Null $|a|$ ist.²
- a ist kleiner als b , wenn a weiter links als b liegt.

Aufgabe 11.1. Machen Sie sich klar, wie man das **Rechnen mit rationalen Zahlen** auf der Zahlengeraden geometrisch deuten kann.

Aus der Definition von $<$ folgen sofort diverse Eigenschaften.

Lemma 11.1

Für alle $a, b, c \in \mathbb{Q}$ gilt $a < b \wedge b < c \Rightarrow a < c$. Ferner gilt immer *genau eine* der drei Aussagen $a = b$, $a < b$ oder $b < a$.

Wir führen noch eine weitere Notation ein, deren Idee in Aufgabe 3.1 bereits gespoilert wurde.

Definition

Für a, b schreiben wir $a \leq b$ als Abkürzung für $a < b \vee a = b$.
 $a \geq b$ ist eine andere Schreibweise für $b \leq a$.

Aus $a < b$ folgt also $a \leq b$, aber aus $a \leq b$ folgt nicht zwangsläufig $a < b$. Man kann nun viele weitere Schlüsse ziehen, die in den folgenden Lemmata zusammengefasst sind.

Lemma 11.2

Für alle $a, b, c \in \mathbb{Q}$ gelten die folgenden Aussagen:

- (i) $a \leq a$
- (ii) $a \leq b \wedge b \leq c \Rightarrow a \leq c$
- (iii) $a \leq b \wedge b \leq a \Rightarrow a = b$
- (iv) $a \leq b \vee b \leq a$

Lemma 11.3

Für alle $a, b, c \in \mathbb{Q}$ gelten die folgenden Aussagen:

- (i) $a < b \Leftrightarrow a + c < b + c$
- (ii) $0 < a \wedge 0 < b \Rightarrow 0 < ab$
- (iii) $a \leq b \Leftrightarrow a + c \leq b + c$
- (iv) $0 \leq a \wedge 0 \leq b \Rightarrow 0 \leq ab$
- (v) $0 < a \Leftrightarrow -a < 0$
- (vi) $0 \leq a \wedge 0 \leq b \Rightarrow 0 \leq a + b$

²Erinnern Sie sich an die Bemerkung zum Absolutbetrag auf Seite 45. Falls Sie sich fragen, wie der Abstand „gemessen“ wird: In der Mathematik gibt man in der Regel keine Einheiten an. Das Stück von 0 bis 1 auf der Zahlengeraden definiert, wie lang die Länge 1 ist.

- (vii) $0 \leq a^2$
- (viii) $0 < 1$
- (ix) $a < b \Leftrightarrow 0 < b - a$
- (x) $a < b \wedge 0 < c \Rightarrow ac < bc$
- (xi) $a < b \wedge c < 0 \Rightarrow bc < ac$
- (xii) $1 < a \wedge 0 < b \Rightarrow b < ab$
- (xiii) $a < 1 \wedge 0 < b \Rightarrow ab < b$
- (xiv) $0 < ab \Leftrightarrow (0 < a \wedge 0 < b) \vee (a < 0 \wedge b < 0)$
- (xv) $0 < a \Leftrightarrow 0 < 1/a$
- (xvi) $1 < a \Leftrightarrow 0 < 1/a < 1$
- (xvii) $0 < a < b \Leftrightarrow 0 < 1/b < 1/a$

Für alle $a, b \in \mathbb{Q}^+$ und alle $m, n \in \mathbb{N}^+$ gilt:

- (i) $a < b \Rightarrow a^n < b^n$
- (ii) $a < b \Rightarrow a^{-n} > b^{-n}$
- (iii) $1 < a \wedge m < n \Rightarrow a^m < a^n$
- (iv) $a < 1 \wedge m < n \Rightarrow a^m > a^n$

Korollar 11.4

Aufgabe 11.2. Die obigen Aussagen müssten Ihnen eigentlich geläufig sein. Trotzdem sollten Sie diese alle einmal durchgehen, sie mit Beispielen nachprüfen und sich jeweils überlegen, was sie geometrisch (also auf der Zahlengeraden) bedeuten.

Aufgabe 11.3. In Lemma 11.3 steht an vielen Stellen eine **Implikation**. In diesen Fällen können Sie davon ausgehen, dass die Aussagen links und rechts vom Implikationspfeil *nicht* äquivalent sind. Beispielsweise ist (ii) keine Äquivalenz, weil aus $0 < ab$ nicht immer $0 < a \wedge 0 < b$ folgt. Überlegen Sie sich für alle Implikationen jeweils Gegenbeispiele, die belegen, dass es sich nicht um Äquivalenzen handelt.

Aufgabe 11.4. a, b, c, d, e und f seien rationale Zahlen mit $a < b$, $0 < c < 1$, $d \geq 0$ und $e < 0$. Welche der folgenden Aussagen folgen aus diesen Angaben? Geben Sie bei den Aussagen, die nicht aus den Voraussetzungen folgen, jeweils ein konkretes Gegenbeispiel an.

- (i) $cd > 0$
- (ii) $a/e > b/e$
- (iii) $a < a + d$
- (iv) $ce < 0$
- (v) $cf < f$
- (vi) $f^2 > 0$
- (vii) $bd \geq 0$
- (viii) $ce \leq 0$
- (ix) $b - e \geq a - e$
- (x) $ae^2 < be^2$
- (xi) $ad < bd$

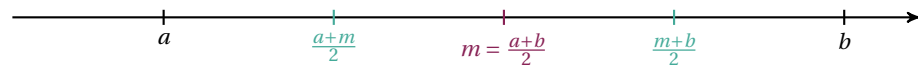
(xii) $d/c \geq d$

Aufgabe 11.5. Welche der folgenden Aussagen sind für alle rationalen Zahlen a , b und c wahr? Geben Sie bei den Aussagen, die nicht wahr sind, jeweils ein konkretes Gegenbeispiel an.

- (i) $a = b \Leftrightarrow a + c = b + c$
- (ii) $a = b \Leftrightarrow ac = bc$
- (iii) $a = b \Leftrightarrow a^2 = b^2$
- (iv) $a = 0 \vee b = 0 \Leftrightarrow ab = 0$
- (v) $a \leq b \Leftrightarrow a + c \leq b + c$
- (vi) $a \leq b \Leftrightarrow ac \leq bc$
- (vii) $a > 0 \wedge b > 0 \Leftrightarrow a + b > 0$
- (viii) $a > 0 \wedge b > 0 \Leftrightarrow ab > 0$

Aufgabe 11.6. Seien a , b , c und d rationale Zahlen. Aus $a = c$ und $b = d$ folgt bekanntlich $a + b = c + d$ und $a - b = c - d$. Aus $a < c$ und $b < d$ folgt jedoch nur $a + b < c + d$, aber nicht zwangsläufig $a - b < c - d$. Geben Sie ein Gegenbeispiel an, also Zahlen, für die $a < c$, $b < d$ aber *nicht* $a - b < c - d$ gilt.

Aus dem bisher Besprochenen folgt eine ganz wesentliche Eigenschaft der rationalen Zahlen. Sind a und b zwei rationale Zahlen mit $a < b$, dann liegt das **arithmetische Mittel** $m = (a + b)/2$ der beiden zwischen a und b , d. h., es gilt $a < m < b$. Man kann nun den Mittelwert von a und m bilden, der zwischen a und m liegt. Analog liegt der Mittelwert von m und b zwischen m und b .



Offenbar kann man diesen Prozess beliebig fortsetzen und erhält ständig neue Zahlen, die alle zwischen a und b liegen. Es gilt also:

Lemma 11.5

Sind a und b rationale Zahlen mit $a < b$, so gibt es unendlich viele rationale Zahlen q mit $a < q < b$.

Das ist eine durchaus unintuitive Eigenschaft, an die man sich erst einmal gewöhnen muss. Wie stark man auch immer an die Zahlengerade „heranzoomt“, man sieht in jedem noch so kleinen Bereich unendlich viele rationale Zahlen. Bis auf die „Namen“ der Zahlen sehen gewissermaßen alle endlichen Stücke der Zahlengeraden unabhängig von ihrer Breite gleich aus.

Definition

Sei M eine Menge von rationalen Zahlen. Eine Zahl $s \in \mathbb{Q}$ heißt **obere** bzw. **untere Schranke** von M (engl. *upper/lower bound*), wenn $s \geq m$ bzw. $s \leq m$ für alle $m \in M$ gilt. Gibt es so eine Zahl, dann sagt man, dass M **nach oben** bzw. **nach unten beschränkt** ist. Ist M nach oben *und* unten beschränkt, so sagt man, dass M **beschränkt** ist.

Die Menge $B_1 = \{r \in \mathbb{Q} : r < 3\}$ ist beispielsweise nach oben, aber nicht nach unten beschränkt. (Daher ist sie nicht beschränkt.) Obere Schranken sind 3, aber auch $7/2$ oder 42. Die Menge $B_2 = \{2, 4/7, -7\}$ ist beschränkt. Obere Schranken sind z. B. 2 oder 23, als untere Schranken könnte man -7 oder -10^{10} nehmen. $\mathbb{N}$ ist nicht beschränkt, aber nach unten durch 0 oder jede negative Zahl beschränkt.

Eine Menge $M \subseteq \mathbb{Q}$ ist genau dann beschränkt, wenn es eine Zahl $s \in \mathbb{Q}$ mit $|m| \leq s$ für alle $m \in M$ gibt.

Lemma 11.6

Aufgabe 11.7. Überlegen Sie sich, was Begriffe wie *obere Schranke* oder *nach unten beschränkt* für die Darstellung auf der Zahlengerade bedeuten.

Aufgabe 11.8. Geben Sie für die folgenden Mengen jeweils an, ob sie nach oben bzw. nach unten beschränkt sind.

$$A_1 = \{m/n : m, n \in \mathbb{Z} \wedge n \neq 0 \wedge m + n \leq 1\,000\,000\}$$

$$A_2 = \mathbb{N} \setminus \mathbb{P}$$

$$A_3 = \{m/n : m, n \in \mathbb{N}^+ \wedge m + n \leq 1\,000\,000\}$$

$$A_4 = \{a \bmod b : a, b \in \mathbb{N}^+ \wedge b \leq 42\}$$

$$A_5 = \{x \in \mathbb{Q}^+ : x^2 \leq x\}$$

$$A_6 = \{x \in \mathbb{Q} : 0 \leq x \leq 42\} \setminus \mathbb{N}$$

$$A_7 = \mathbb{N} \setminus \{x \in \mathbb{Q} : 0 \leq x \leq 42\}$$

$$A_8 = \mathbb{N} \cap \{x \in \mathbb{Q} : 0 \leq x < 42\}$$

$$A_9 = \{a^{p-1} \bmod p : a \in \mathbb{N} \wedge p \in \mathbb{P}\}$$

$$A_{10} = \{x \in \mathbb{Q}^+ : x^3 \leq x\}$$

$$A_{11} = \{x \in \mathbb{Q} : x^3 \leq x\}$$

$$A_{12} = \{m/n : m, n \in \mathbb{N} \wedge m < n\}$$

Aufgabe 11.9. Sei $M \subseteq \mathbb{Q}$ eine Menge, die nach unten durch a und nach oben durch b beschränkt ist. Wie kommt man mit diesen Angaben auf die Zahl s aus Lemma 11.6?

Sei M eine Menge von rationalen Zahlen. Eine obere Schranke von M , die ein Element von M ist, nennt man **Maximum** oder **größtes Element** von M und schreibt für dieses Element $\max M$. Eine untere Schranke von M , die ein Element von M ist, nennt man **Minimum** oder **kleinstes Element** von M und schreibt für dieses Element $\min M$.

Definition

Wenn es ein Maximum gibt, so ist dieses offenbar eindeutig bestimmt. Man spricht also von *dem* Maximum. (Und das gilt natürlich auch für *das* Minimum.)

Aufgabe 11.10. Aus der Definition folgt, dass eine Menge, die ein Maximum hat, durch dieses Maximum nach oben beschränkt sein muss. Geben Sie eine Menge an, die nach oben beschränkt ist, die aber kein Maximum hat.

Aufgabe 11.10 hat gezeigt, dass es Mengen ohne Maximum gibt, und natürlich gibt es auch welche ohne Minimum. Unter gewissen Bedingungen gibt es jedoch immer solche Elemente.

Lemma 11.7

Jede nichtleere³ endliche Teilmenge von $\mathbb{Q}$ hat ein Minimum und ein Maximum und ist daher beschränkt.

Die Aussage folgt gewissermaßen daraus, dass man die Elemente der Reihe nach durchgehen könnte. (Siehe dazu den folgenden [Abschnitt über Ungleichungen im Computer](#).)

Schließlich noch eine grundlegende Eigenschaft der natürlichen Zahlen, die wir mit unseren Mitteln nicht beweisen können, sondern einfach glauben müssen.

Satz 11.8**Wohlordnung der natürlichen Zahlen**

Jede nichtleere Menge von natürlichen Zahlen hat ein Minimum.

Aufgabe 11.11. Beachten Sie, dass Satz 11.8 nicht für die ganzen bzw. die rationalen Zahlen gilt. Geben Sie Beispiele für nichtleere Teilmengen von $\mathbb{Z}$ bzw. $\mathbb{Q}$ an, die kein Minimum haben.

Aufgabe 11.12. Satz 11.8 gilt auch nicht mehr, wenn man das Minimum durch das Maximum ersetzt. Geben Sie ein Beispiel für eine nichtleere Teilmenge von $\mathbb{N}$ an, die kein Maximum hat.

★**Aufgabe 11.13.** Die Wohlordnung der natürlichen Zahlen spielt oft auch in der Informatik eine Rolle; unter anderem, wenn bewiesen werden soll, warum ein Programm funktioniert. Können Sie z. B. begründen, warum der [euklidische Algorithmus](#) immer terminiert?

Bemerkungen

Man beachte, dass die Zahl null weder positiv noch negativ ist. Siehe dazu auch die Bemerkung auf Seite 39.

Bei der Frage nach Schranken wird von machen Studierenden ∞ ins Spiel gebracht. Daher sei hier noch einmal daran erinnert, dass ∞ *keine Zahl* ist und daher auch keine Schranke sein kann. (Siehe die Lösung von Aufgabe 4.19.)

In den [Restklassenkörpern](#) der Form $\mathbb{Z}_p$ kann man zwar rechnen wie in $\mathbb{Q}$, es gibt aber keine vergleichbare Ordnung, die sich mit den Rechenoperationen im Sinne von Lemma 11.3 verträgt. Es ergibt also beispielsweise in $\mathbb{Z}_5$ keinen Sinn, zu sagen, dass 2 „kleiner“ als 4 sei. Und auch keine andere Sortierung würde den entsprechenden Zweck erfüllen können.⁴

nichtleer

³Wie Sie sich wohl schon gedacht haben, bezeichnet man in der Mathematik Mengen als *nichtleer*, wenn sie mindestens ein Element enthalten, also nicht die leere Menge sind.

⁴Warum das nicht geht, wird im Abschnitt über [geordnete Körper](#) im [alten Skript](#) erklärt.

★ Ungleichungen im Computer

PYTHON verwendet wie viele andere Programmiersprachen für $<$ und $\leq$ die Zeichen $<$ und $<=$. Da Computer aber nur vergleichsweise wenige Zahlen exakt darstellen können (mehr dazu [im nächsten Abschnitt](#)), liegen zwischen zwei [Fließkommazahlen](#) nicht unendlich viele andere. PYTHON hat auch Funktionen für Maximum und Minimum:

```
B2 = {2, 4/7, -17}
[f(B2) for f in [max, min]]
```

Die hätte man aber auch selbst schreiben können; für Mengen z. B. so:

```
def Max(M):
    m = next(iter(M))    # siehe Fußnote 37 auf Seite 53
    while len(M)>0:
        k = next(iter(M))
        if k>m: m = k
        M = M-{k}         # nicht remove, um M zu "schützen"
    return m
```

Dieses Programm, das aus didaktischen Gründen keine for-Schleife verwendet, sollte nebenbei begründen, warum in Lemma 11.7 die Existenz eines Maximums für endliche Mengen garantiert werden kann: weil der Algorithmus für solche Mengen auf jeden Fall terminiert, da $|M|$ in jedem Schritt um 1 verringert wird. (Eine weitere Anwendung von Satz 11.8.)

12. Die reellen Zahlen

Die reellen Zahlen werden zwar in der Schule ständig verwendet, jedoch selten ausführlich hinterfragt. Dazu fehlt erstens die Zeit und zweitens wird die Materie auch schnell kompliziert. Wir werden dem Thema auch nicht auf den Grund gehen, wollen aber zumindest herausarbeiten, welche „Lücke“ in der Zahlenwelt die reellen Zahlen schließen und warum sie nicht einfach zu handhaben sind.

Wenn Sie mit einem Taschenrechner oder einer Programmiersprache die Quadratwurzel von $3/5$ berechnen, werden Sie eine Antwort wie

$$0.7745966692414834 \quad (12.1)$$

erhalten. Das ist jedoch *nicht* die Quadratwurzel von $3/5$! Wieso? Damit die Sache nicht zu unübersichtlich wird, nehmen wir an, die Antwort wäre 0.775 gewesen. Das ist der Bruch $775/1000$ bzw. in gekürzter Form $31/40$. Wäre das die Quadratwurzel von $3/5$, so würde die Gleichung

$$\left(\frac{31}{40}\right)^2 = \frac{31 \cdot 31}{40 \cdot 40} = \frac{961}{1600} \stackrel{?}{=} \frac{3}{5} \quad (12.2)$$

gelten. Man „sieht“, dass das nicht stimmt. Aber kann man so etwas auch bei großen Zählern und Nennern wie beispielsweise (12.1) sofort erkennen? Man kann!

Weil $31/40$ in gekürzter Form vorliegt, sind Zähler und Nenner **teilerfremd**. Offensichtlich sind dann auch Zähler und Nenner von $31^2/40^2$ teilerfremd. Weil aber auch $3/5$ ein gekürzter Bruch ist, stehen links und rechts vom letzten Gleichheitszeichen in (12.2) definitiv zwei verschiedene Brüche. Anders ausgedrückt: Wenn p/q ein gekürzter Bruch ist, dann ist p^2/q^2 die gekürzte Darstellung seines Quadrats. Noch allgemeiner formuliert:

Satz 12.1**Satz des Theaitetos**

Ist r eine nichtnegative rationale Zahl und $n \in \mathbb{N}^+$, so hat die Gleichung $x^n = r$ dann und *nur* dann eine *rationale* Lösung, wenn es natürliche Zahlen p und q gibt, so dass p^n/q^n die gekürzte Darstellung von r ist.

Gibt es eine Lösung, so ist diese trivialerweise p/q . Interessanter ist die Aussage, wann es *keine* Lösung gibt. Beispielsweise kann $x^2 = 2$ keine rationale Lösung haben, denn die gekürzte Darstellung von 2 ist $2/1$ und der Zähler 2 ist keine Quadratzahl.⁵ Ebenso haben Gleichungen wie $x^2 = 3$ oder $x^2 = 6$ keine rationale Lösung. Und auch $x^5 = 5/32$ hat beispielsweise keine, denn der Zähler 5 ist offensichtlich nicht die fünfte Potenz einer natürlichen Zahl.

Aufgabe 12.1. Welche der folgenden Gleichungen haben rationale Lösungen?

$x^2 = 13$

$x^2 = 121$

$x^2 = \frac{45}{80}$

$x^5 = \frac{1}{32}$

$x^3 = 27$

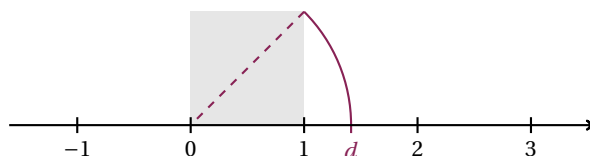
$x^3 = 29$

$x^4 = \frac{16}{125}$

$x^3 = \frac{3}{81}$

Aufgabe 12.2. In der Lösung von Aufgabe 12.1 steht, dass eine rationale Lösung von $x^3 = 29$ eine natürliche Zahl sein müsste. Wieso stimmt das?

Im **letzten Abschnitt** stand, dass jeder rationalen Zahl ein Punkt auf der Zahlengerade entspricht. Umgekehrt entspricht aber nicht jedem Punkt auf der Zahlengerade eine rationale Zahl, wie man sich durch eine einfache geometrische Konstruktion überlegen kann. Man zeichnet ein Quadrat, dessen eine Seite das Stück von 0 bis 1 ist. Nach dem **Satz des Pythagoras** gilt für die Länge d der Diagonale dieses Quadrats $d^2 = 1^2 + 1^2 = 2$. Überträgt man diese Diagonale wie in der folgenden Skizze mit einem Zirkel auf die Zahlengerade, so trifft man offenbar einen Punkt, der nach dem **Satz des Theaitetos** keiner rationalen Zahl entsprechen kann.



Diese „fehlenden“ Punkte sind ein Indiz dafür, dass man weitere Zahlen braucht. Dass der betreffende Satz nach dem griechischen Mathematiker **Theaitetos** benannt ist, der vor zweieinhalb Jahrtausenden lebte, zeigt, wie lange dieses Manko

Quadratzahl

⁵Als *Quadratzahlen* bezeichnet man die Quadrate natürlicher Zahlen. Die Menge der Quadratzahlen ist also $\{n^2 : n \in \mathbb{N}\} = \{0, 1, 4, 9, 16, 25, \dots\}$.

bereits bekannt ist. Erstaunlicherweise war die Lösung des Problems, die inzwischen allgemeiner Usus in der Mathematik ist, bis zur Mitte des 20. Jahrhunderts aus philosophischen Gründen umstritten. (Und eine kleine Minderheit ist nach wie vor der Meinung, dass die besagte Lösung gar keine ist. Mehr dazu in dem Video [Das Rätsel des Kontinuums](#).)

Wir werden die Geschichte der Mathematik jedoch ignorieren und uns stattdessen als Vorbereitung auf die Einführung der reellen Zahlen mit der Dezimaldarstellung von Zahlen, wie man sie beispielsweise in (12.1) sieht, beschäftigen.⁶

Wir verwenden ein sogenanntes [Stellenwertsystem](#), mit dem man – anders als z. B. mit der [römischen Zahlschrift](#) – mit ganz wenigen Ziffern beliebig große Zahlen schreiben kann. Der „Trick“ dabei ist, dass der Wert einer Ziffer von ihrer Position (*Stelle*) in der Zahl abhängt. Der Wert der Ziffer 4 in 941 ist beispielsweise 40, während der Wert derselben Ziffer in 419 die Zahl 400 ist. Sie kennen das natürlich im Prinzip alles, weil Sie Zahlen schon seit Ihrer Kindheit so geschrieben haben. Die Stellen entsprechen Zehnerpotenzen, die rechts mit 10^0 – also 1 – anfangen und nach links jeweils um den Faktor 10 größer werden:⁷

$$9419 = 9 \cdot 1000 + 4 \cdot 100 + 1 \cdot 10 + 9 \cdot 1 = 9 \cdot 10^3 + 4 \cdot 10^2 + 1 \cdot 10^1 + 9 \cdot 10^0.$$

Dem im 16. Jahrhundert tätigen flämischen Mathematiker [Simon Stevin](#) haben wir eine logische Weiterentwicklung dieses Systems zu verdanken: die [Nachkommastellen](#). Man schreibt dabei rechts hinter einem [Trennzeichen](#) weiter und die Werte entstehen nun durch Zehnerpotenzen mit negativen Exponenten:

$$94.19 = 9 \cdot 10 + 4 \cdot 1 + 1 \cdot \frac{1}{10} + 9 \cdot \frac{1}{100} = 9 \cdot 10^1 + 4 \cdot 10^0 + 1 \cdot 10^{-1} + 9 \cdot 10^{-2} \quad (12.3)$$

In diesem Skript wird übrigens als Dezimaltrennzeichen grundsätzlich ein Punkt verwendet, obwohl in Deutschland das Komma üblicher ist. Diese Entscheidung beruht auf ganz pragmatischen Gründen. Erstens sind Dezimalpunkte in den meisten Programmiersprachen üblich und zweitens kann es dann keine Verwechslungsgefahr mit Kommata in Mengen oder Tupeln geben. Wir werden trotzdem von *Nachkommastellen* sprechen. Offensichtlich kann man sich beim Studium dieser Stellen auf Zahlen zwischen 0 und 1 konzentrieren, weil der ganzzahlige Anteil vor dem Komma dann uninteressant ist.

⁶Diese Darstellung verdanken wir übrigens der indischen Mathematik. Dort ist sie vor circa 1500 Jahren entstanden und dann über das arabische Reich schließlich nach Europa gekommen, wo sie anfangs nur sehr zögerlich angenommen wurde. Das wird in der dreiteiligen [arte-Dokumentation *Die Odyssee der Zahlen*](#) sehr schön erzählt.

⁷Dass die Basis des Systems die Zahl 10 ist, ist gewissermaßen ein historischer „Zufall“. Wahrscheinlich liegt es daran, dass Menschen zehn Finger haben. Mathematisch ist die Zehn keine besondere Zahl und man könnte statt ihrer auch jede andere nehmen. Sie wissen wahrscheinlich, dass Computer intern mit [einem Stellenwertsystem zur Basis 2](#) rechnen.

Fangen wir mit der einfachsten Variante an. Zahlen, die nur endlich viele Nachkommastellen haben,⁸ kann man offenbar immer als Brüche darstellen, deren Nenner eine Zehnerpotenz ist:

$$\begin{aligned} 0.9419 &= \frac{9}{10} + \frac{4}{100} + \frac{1}{1000} + \frac{9}{10000} \\ &= \frac{9000}{10000} + \frac{400}{10000} + \frac{10}{10000} + \frac{9}{10000} = \frac{9419}{10000} \\ 0.445 &= \frac{455}{1000} = \frac{91}{200} \\ 0.0306 &= \frac{306}{10000} = \frac{153}{5000} \end{aligned}$$

Wie man an den Beispielen sieht, kann man evtl. noch kürzen, so dass der Nenner des gekürzten Bruchs dann keine Zehnerpotenz mehr ist.

Umgekehrt kann man für manche Brüche durch einfaches Erweitern ihre Dezimaldarstellung ermitteln:

$$\begin{aligned} \frac{3}{4} &= \frac{3 \cdot 25}{4 \cdot 25} = \frac{75}{100} = 0.75 \\ \frac{17}{5000} &= \frac{17 \cdot 2}{5000 \cdot 2} = \frac{34}{10000} = 0.0034 \end{aligned}$$

Aufgabe 12.3. Für den Fall, dass Ihnen das nicht ohnehin klar ist, sollten Sie es üben, indem Sie ein paar ausgedachte Zahlen von Hand umwandeln und dann mit einem Taschenrechner Ihre Ergebnisse überprüfen. Beachten Sie dabei insbesondere auch die Rolle führende Nullen wie bei den Beispielen 0.0306 und 17/5000 oben. Was passiert mit den Dezimaldarstellungen, wenn man mit Zehnerpotenzen multipliziert oder durch sie dividiert?

Was folgt daraus? Zahlen, die sich *nicht* in diese Form bringen lassen, kann man nicht mit endlich vielen Nachkommastellen darstellen. Diese Zahlen fallen in zwei Kategorien:

- Brüche, deren Nenner in gekürzter Form außer 2 und 5 noch andere **Primfaktoren** haben, denn durch Erweitern wird man diese Faktoren nicht „loswerden“ und daher im Nenner keine Zehnerpotenz erhalten.
- Zahlen, die gar keine Brüche sind: die sogenannten *irrationalen* Zahlen. Zu denen kommen wir **demnächst**.

Wie erhält man generell für beliebige Brüche die Dezimaldarstellung? Das kann man durch **schriftliche Division** erreichen. Analysieren wir das am Beispiel $3/7$, wobei wir $3.0000\dots$ statt 3 schreiben und uns vorstellen, dass wir beliebig viele Nullen „anhängen“ können:

⁸Wenn von endlich vielen Nachkommastellen gesprochen wird, dann ist das natürlich immer so gemeint, dass es sich um endlich viele handelt, die von null verschieden sind. Theoretisch kann man 0.5 ja auch als 0.50, 0.500000 oder $0.5\overline{0}$ schreiben.

$$\begin{array}{r}
 3.0000000 \dots : 7 = 0.4285714 \dots \\
 \underline{-2.8} \\
 20 \\
 \underline{-14} \\
 60 \\
 \underline{-56} \\
 40 \\
 \underline{-35} \\
 50 \\
 \underline{-49} \\
 10 \\
 \underline{-7} \\
 30 \\
 \underline{-28} \\
 2
 \end{array}$$

Die erste Nachkommastelle 4 haben wir erhalten, indem wir im Prinzip 30 durch 7 dividiert haben. Wenn wir uns nicht verrechnet haben, ist $30/7$ also mindestens 4, evtl. etwas mehr (wenn es einen Rest gibt), aber weniger als 5 (denn sonst hätten wir doch einen Fehler gemacht). Da wir eine Null „geborgt“ haben, haben wir eigentlich herausbekommen, dass $3/7$ mindestens 0.4 und auf jeden Fall weniger als 0.5 ist. In diesem Sinne ist 4 die „richtige“ erste Nachkommastelle. Analog sagt uns die zweite Nachkommastelle, dass $3/7$ mindestens 0.42 und weniger als 0.43 ist. Und so geht es weiter.

Weil dieses Vorgehen offenbar korrekt funktioniert, muss es bei Brüchen, die sich mit endlich vielen Nachkommastellen darstellen lassen, abbrechen. Das heißt, dass als **Rest** (das sind im Beispiel die blauen Zahlen) irgendwann null herauskommt.

Was passiert bei den anderen Brüchen? Im Beispiel kommt es zu einer Wiederholung. Ganz unten kommt der Rest 2 heraus, der weiter oben schon einmal vorkam. So eine Wiederholung *musste* sich ergeben, weil bei der Division durch 7 außer der Null nur sechs verschiedene Reste möglich sind. Und offenbar gilt das immer. Bei der Division durch 23 muss es spätestens nach 22 Divisionen zu einer Wiederholung gekommen sein und so weiter.⁹

Wir können damit festhalten, dass es in der Dezimaldarstellung von rationalen Zahlen immer zu **periodischen Wiederholungen** kommen wird.¹⁰ Die entsprechende Schreibweise mit einem **Überstrich** sollte Ihnen aus der Schule geläufig sein; für $3/7$ ist es beispielsweise $0.\overline{428571}$.

⁹Man kann das noch verfeinern und genauer vorhersagen, wann es zum ersten Mal passiert. Dazu gibt es ein Video von mir, das den Titel *Warum ist $1/7$ interessanter als $1/13$?* trägt.

¹⁰Damit der Satz auch für abbrechende Darstellungen stimmt, muss man Zahlen wie 0.32 als $0.32\overline{0}$ oder $0.31\overline{9}$ interpretieren.

Aber wie genau sind „unendlich viele Nachkommastellen“ zu verstehen? Für die Zahl $a = 0.\overline{27}$ gilt beispielsweise

$$\begin{aligned} 0.2 &\leq a \leq 0.3, \\ 0.27 &\leq a \leq 0.28, \\ 0.272 &\leq a \leq 0.273, \\ 0.2727 &\leq a \leq 0.2728 \end{aligned} \tag{12.4}$$

und immer so weiter. Man kann das so interpretieren, dass die ersten n Nachkommastellen so einer Zahl diese niemals genau spezifizieren, ihre Lage aber zwischen zwei benachbarten Zahlen mit dieser Anzahl von Nachkommastellen eingrenzen. Multipliziert man die obigen Ungleichungen mit 100, so erhält man

$$\begin{aligned} 20 &\leq 100a \leq 30, \\ 27 &\leq 100a \leq 28, \\ 27.2 &\leq 100a \leq 27.3, \\ 27.27 &\leq 100a \leq 27.28, \\ 27.272 &\leq 100a \leq 27.273, \quad \dots \end{aligned}$$

und damit

$$\begin{aligned} 0.2 &\leq 100a - 27 \leq 0.3, \\ 0.27 &\leq 100a - 27 \leq 0.28, \\ 0.272 &\leq 100a - 27 \leq 0.273 \end{aligned}$$

et cetera. Das ist dieselbe Kette von Ungleichungen wie (12.4). Daraus kann man schließen, dass $100a - 27$ und a dieselbe Zahl sein müssen.¹¹ Löst man $100a - 27 = a$ nach a auf, so erhält man $a = 27/99 = 3/11$. Der „Trick“ dabei war, dass wir mit 10^2 multipliziert haben, um die zweistellige Periode vor das Komma zu schieben. Die Zahl 99 ergab sich, weil wir 1 von 10^2 subtrahiert haben. Bei einer fünfstelligen Periode wie $0.\overline{02439}$ würden wir mit 10^5 multiplizieren und kämen auf den Nenner $10^5 - 1 = 99999$, also

$$0.\overline{02439} = \frac{2439}{99999} = \frac{1}{41}.$$

Aufgabe 12.4. Geben Sie $0.\overline{24}$ und $0.\overline{370}$ jeweils als gekürzten Bruch an.

Aufgabe 12.5. Wenden Sie diese Methode auf $0.\overline{9}$ an.

Aufgabe 12.6. Überprüfen Sie durch schriftliche Division, dass man für den Bruch $3/11$ tatsächlich die Dezimaldarstellung $0.\overline{27}$ erhält.

¹¹Wenn man es ganz genau nimmt, kann man das erst nach einer sauberen Definition der **reellen Zahlen** schließen ...

Daraus kann man eine allgemeine Methode machen, eine Zahl der Form

$$0.k_1 \dots k_m \overline{p_1 \dots p_n}$$

in einen Bruch umzurechnen. Der vordere Teil ist einfach die Zahl $0.k_1 \dots k_m$ und daraus erhält man den Bruch mit dem Zähler $k_1 \dots k_m$ und dem Nenner 10^m . Der periodische Teil hinten ist $10^{-m} \cdot 0.\overline{p_1 \dots p_n}$ und $0.\overline{p_1 \dots p_n}$ wiederum ist der Bruch mit dem Zähler $p_1 \dots p_n$ und dem Nenner $10^n - 1$. Die beiden Teilergebnisse müssen dann nur noch addiert werden.

Am Beispiel $0.0012\overline{017}$ ausführlich vorgerechnet:

- Hier haben wir $m = 4$, $k_1 \dots k_4 = 0012$, $n = 3$ und $p_1 \dots p_3 = 017$.
- Der nichtperiodische Teil ist $12/10^4 = 3/2500$.
- $0.\overline{017}$ ist $17/999$.
- $0.0000\overline{017}$ ist daher $10^{-4} \cdot 17/999 = 17/9990000$.
- Insgesamt ergibt sich $3/2500 + 17/9990000 = 2401/1998000$.

Aufgabe 12.7. Wenn Sie es noch einmal selbst probieren wollen: Geben Sie die Dezimalzahl $0.7\overline{63}$ als gekürzten Bruch an.

Aufgabe 12.8. Wenn Sie den Umgang mit unendlichen Dezimaldarstellungen üben wollen, können Sie beispielsweise [Wolfram Alpha](#) verwenden, um weitere Aufgaben zu generieren. Gibt man dort etwa

decimal digits of 49/660

ein, dann erhält man als Antwort $\{\{7, \{4, 2\}\}, -1\}$. Dabei steht der erste Teil für den nichtperiodischen Anfang, der zweite für die Periode. Hier geht es somit um die Ziffernfolge $74\overline{2}$. Die letzte Zahl gibt an, wie viele der Ziffern der Folge links vom Dezimaltrennzeichen stehen. Weil die Zahl in diesem Fall negativ ist, fängt die Sequenz erst mit einer Stelle „Verspätung“ an: es kommt also $0.07\overline{42}$ heraus.

Man kann sich nun zumindest vorstellen, dass es auch Zahlen gibt, die unendlich viele Nachkommastellen haben, ohne dass es jemals zu periodischen Wiederholungen kommt. Das könnte beispielsweise

$$42.12345678910111213141516171819202122\dots \quad (12.5)$$

sein. Nimmt man solche Zahlen zu den rationalen hinzu, so kommt man zu der folgenden Definition, die mathematisch nicht ganz präzise ist, für unsere Zwecke jedoch ausreicht.

Durch jede Folge von endlich oder unendlich vielen Dezimalziffern mit optionalem Vorzeichen und optionalem Dezimaltrennzeichen wird eindeutig eine **reelle Zahl** bestimmt und für jede reelle Zahl gibt es so eine Darstellung. Für die Menge aller reellen Zahlen schreibt man $\mathbb{R}$, für die Menge der positiven (bzw. negativen) reellen Zahlen schreibt man $\mathbb{R}^+$, bzw. $\mathbb{R}^-$. Zahlen, die reell, aber nicht rational sind, nennt man **irrational**.

Definition

Um die Definition nicht zu überfrachten, wurden ein paar Details weggelassen, die hoffentlich klar sind, die aber vorsichtshalber hier erwähnt werden:

Mit dem *Vorzeichen*, das *vor* der Ziffernfolge steht, ist natürlich eines der Symbole + oder – gemeint, wobei + wie üblich der "**Default-Wert**" ist, falls kein Vorzeichen angegeben wurde. Das Vorzeichen entscheidet darüber, ob eine (von null verschiedene) reelle Zahl positiv oder negativ ist.

Handelt es sich um unendlich viele Ziffern, so ist das Dezimaltrennzeichen zwingend erforderlich und links von ihm dürfen nur endlich viele Ziffern stehen.¹² Steht links vom Dezimaltrennzeichen gar keine Ziffer, so wird das wie im englischen Sprachraum und vielen Programmiersprachen so interpretiert, als stünde dort eine Null.

Die Definition ist so gemeint, dass wir eine Zahl wie die in (12.5) analog zu (12.4) so interpretieren, dass diese Zahl beispielsweise zwischen 42.12345 und 42.12346, zwischen 42.123456789101 und 42.123456789102 und so weiter liegt. Die Folge der Nachkommastellen einer reellen Zahl r engt also den Bereich, in dem r liegt, immer weiter ein. (Damit haben wir dann auch die fehlende Begründung, die in Fußnote 11 angesprochen wurde, nachgereicht.)

Das hat unter anderem die folgenden Konsequenzen:

- Jede rationale Zahl ist reell, weil wir bereits gesehen haben, wie man die Nachkommastellen von Brüchen berechnet. Es gilt also $\mathbb{Q} \subseteq \mathbb{R}$.
- Es gilt sogar $\mathbb{Q} \subset \mathbb{R}$, was einfach nur bedeutet, dass es irrationale Zahlen gibt. Die Zahl aus (12.5) ist ein Beispiel für eine irrationale Zahl. Wir werden **bald sehen**, dass Gleichungen wie $x^2 = 2$ reelle Lösungen haben. Da diese nach dem **Satz des Theaitetos** keine Brüche sein können, handelt es sich um irrationale Zahlen. Eine weitere bekannte irrationale Zahl ist π .¹³
- Eine Zahl ist genau dann rational, wenn die Ziffern in der Folge ihrer Nachkommastellen sich ab einer gewissen Stelle periodisch wiederholen. Für Zahlen, die sich in gekürzter Form als Bruch $p/q \neq 0$ darstellen lassen, wobei q außer 2 und 5 keine weiteren Primfaktoren hat, gibt es zwei verschiedene Dezimaldarstellungen. (Siehe Fußnote 10.) Für alle anderen Zahlen ist die Dezimaldarstellung eindeutig bestimmt.

Wenn wir in Zukunft von der n -ten Nachkommastelle einer Zahl sprechen, dann ist bei den Zahlen, für die es zwei Darstellungen gibt, immer die Darstellung mit unendlich vielen Nullen gemeint. Die dritte Nachkommastelle von $a = 21/50$ ist beispielsweise 0 und *nicht* 9, obwohl man a sowohl als $0.42\overline{0}$ als auch als $0.41\overline{9}$ schreiben könnte.

- Wie in (12.3) kann man sich auch bei unendlich vielen Nachkommastellen eine Summe vorstellen, zum Beispiel

$$0.\overline{27} = \frac{2}{10^1} + \frac{7}{10^2} + \frac{2}{10^3} + \frac{7}{10^4} + \frac{2}{10^5} + \frac{7}{10^6} + \dots$$

Für solche Summen mit unendlich vielen Summanden gibt es eine präzise mathematische Beschreibung, zu der wir **später** kommen.

¹²Es gibt also keine „unendlich großen“ oder „unendlich kleinen“ reellen Zahlen.

¹³Das wird in dem Video *Die Kreiszahl Pi ist irrational* so bewiesen, dass man es mit guten Schulkenntnissen verstehen sollte.

Die irrationalen Zahlen werden gewissermaßen zu den rationalen Zahlen hinzugefügt, um „Lücken“ zu schließen. Es gibt dann keine „**fehlenden Punkte**“ mehr: Jede reelle Zahl entspricht einem Punkt auf der Zahlengeraden, aber umgekehrt entspricht auch jedem Punkt auf der Zahlengeraden eine reelle Zahl.

Irrationale Zahlen kann man im Gegensatz zu Brüchen nicht einfach so hinschreiben. Man kann ihnen lediglich „Namen“ wie π oder e geben. Allerdings kann man mit ihnen wie gewohnt rechnen:

Alle Regeln der Arithmetik aus Abschnitt 5 gelten auch für die reellen Zahlen. Das gilt ebenfalls für die Eigenschaften aus Abschnitt 11. Insbesondere gibt es zwischen je zwei verschiedenen reellen Zahlen unendlich viele rationale und unendlich viele irrationale Zahlen.

Satz 12.2

Dass es zwischen zwei reellen Zahlen mindestens eine weitere (und damit unendlich viele) geben muss, kann man sich einfach dadurch klarmachen, dass die beiden Zahlen sich ab einer bestimmten Nachkommastelle unterscheiden müssen (sonst wären sie gleich). Das kann man ausnutzen, um eine Zahl zwischen die beiden zu „schummeln“. Beispielsweise unterscheiden sich der Bruch $337/8 = 42.125$ und die irrationale Zahl aus (12.5) in der dritten Nachkommastelle, die einmal 5 und einmal 3 ist. Zwischen ihnen liegen Zahlen wie 42.124, 42.1235, 42.1245 und so weiter. (Das sind zwar alles rationale Zahlen, aber man kann sich ohne Probleme auch irrationale Zahlen zwischen den Ausgangszahlen „basteln“, indem man z. B. zu 42.124 unendlich viele Nachkommastellen ohne periodische Wiederholung hinzufügt.)

Aufgabe 12.9. Wie erwähnt ist π irrational. Begründen Sie mithilfe von Satz 5.5, dass $2 + \pi$ und 2π ebenfalls irrational sind.

Aufgabe 12.10. Kann die Summe zwei irrationaler Zahlen rational sein? Und was ist mit dem Produkt?

Für Mengen von reellen Zahlen, die „zwischen“ Punkten auf der Zahlengeraden liegen, führen wir nun offiziell eine Schreibweise und einen Begriff ein. Beide werden typischerweise bereits in der Schule verwendet.

Für $a, b \in \mathbb{R}$ definieren wir die folgenden Mengen:

$$\begin{array}{ll} [a, b] = \{x \in \mathbb{R} : a \leq x \leq b\} & [a, b) = \{x \in \mathbb{R} : a \leq x < b\} \\ (a, b] = \{x \in \mathbb{R} : a < x \leq b\} & (a, b) = \{x \in \mathbb{R} : a < x < b\} \\ [a, \infty) = \{x \in \mathbb{R} : a \leq x\} & (a, \infty) = \{x \in \mathbb{R} : a < x\} \\ (-\infty, b] = \{x \in \mathbb{R} : x \leq b\} & (-\infty, b) = \{x \in \mathbb{R} : x < b\} \\ (-\infty, \infty) = \mathbb{R} & \end{array}$$

Mengen, die man so darstellen kann, nennt man **Intervalle**. Intervalle, die mehr als ein Element enthalten, nennt man **echte Intervalle**.

Definition

Aufgabe 12.11. Wie sehen nach dieser Definition die Mengen $[42, 42]$ und $[42, 42)$ aus?

Anschaulich sind Intervalle zusammenhängende Abschnitte der Zahlengeraden, also Mengen von reellen Zahlen ohne „Lücken“: Enthält ein Intervall zwei Zahlen, so enthält es auch alle Zahlen zwischen diesen beiden. Das schließt allerdings Extremfälle wie in Aufgabe 12.11 nicht aus. Insgesamt erhält man damit die folgende Aussage.

Lemma 12.3

Alle echten Intervalle enthalten unendlich viele reelle Zahlen. Unechte Intervalle haben entweder genau ein Element oder gar keins.

Aufgabe 12.12. Geben Sie $\mathbb{R}^+$ und $\mathbb{R}^-$ in Intervallschreibweise an.

Aufgabe 12.13. Wie viele Elemente hat die Menge $[1, 5]$?

Aufgabe 12.14. Der Durchschnitt von zwei Intervallen ist immer ein Intervall (evtl. aber kein echtes). Geben Sie zur Übung $[1/2, 5) \cap (3, 7]$, $(1/2, 5] \cap [3, 7)$, $[1/2, \infty) \cap (3, 7]$, $[1/2, 3] \cap [3, 7]$ und $[1/2, 3) \cap [3, 7]$ an.

Aufgabe 12.15. Geben Sie zwei Intervalle an, deren Vereinigung ein Intervall ist. Geben Sie dann zwei Intervalle an, deren Vereinigung *kein* Intervall ist.

Aufgabe 12.16. Ist $I_1 \setminus I_2$ immer ein Intervall, wenn I_1 und I_2 Intervalle sind?

Aufgabe 12.17. Welche der folgenden Mengen sind echte Intervalle?

$$[40, 42] \setminus [40, 41)$$

$$[40, 41] \setminus \mathbb{N}$$

$$[40, 42] \setminus [40, 42)$$

$$[40, 41] \setminus \mathbb{Q}$$

$$[40, 41] \cap [41, 43]$$

$$(40, 42) \setminus \mathbb{N}$$

Durch die Einführung der reellen Zahlen gibt es gleichsam ein „Happy End“ für das Problem des Theaitetos:

Satz 12.4

Ist y eine nichtnegative reelle Zahl und $n \in \mathbb{N}^+$, dann gibt es eine eindeutig bestimmte nichtnegative reelle Zahl x mit $x^n = y$.

Man kann diese Zahl erhalten, indem man sich ihr mit rationalen Zahlen immer weiter annähert. Durch Probieren ergibt sich z. B.

$$1.4^2 < 2 < 1.5^2,$$

$$1.41^2 < 2 < 1.42^2,$$

$$1.414^2 < 2 < 1.415^2,$$

$$1.4142^2 < 2 < 1.4143^2$$

und so weiter.¹⁴ Damit ist klar, dass die ersten vier Nachkommastellen der Lösung der Gleichung $x^2 = 2$ die Ziffern 4, 1, 4 und 2 sein müssen. Hinter der Definition der reellen Zahlen steht die Idee, dass man diesen Prozess bis ins Unendliche ausdehnt und dadurch einen Wert festlegt. Er bekommt den Namen $\sqrt{2}$:

Die eindeutig bestimmte Zahl x aus Satz 12.4 nennt man die **n -te Wurzel** von y und schreibt dafür $\sqrt[n]{y}$. $\sqrt[2]{y}$ nennt man auch die **Quadratwurzel** von y und schreibt dafür kurz $\sqrt{y}$.

Definition

Es ist wichtig, dass der Ausdruck $\sqrt[n]{y}$ durch diese Definition eindeutig bestimmt ist. $\sqrt{4}$ ist z. B. 2 und *nicht* -2 oder so etwas wie „ ± 2 “ (obwohl das Quadrat von -2 natürlich 4 ist). Außerdem kann man die Gleichung $x^n = y$ zwar auch für negative y lösen, wenn n ungerade ist (beispielsweise gilt $(-3)^3 = -27$), aber man verwendet die Wurzelschreibweise in der Regel nur für nichtnegative y .

Für $y \in \mathbb{R}$ mit $y \geq 0$, $p \in \mathbb{Z}$ und $q \in \mathbb{N}^+$ definieren wir $y^{\frac{p}{q}}$ als $\sqrt[q]{y^p}$.

Definition

Damit ist ein Ausdruck der Form y^r nun für beliebige rationale Exponenten r definiert. Und wegen $\sqrt[n]{y} = y^{1/n}$ widerspricht diese Definition nicht [der vorherigen](#) für ganzzahlige Exponenten, sondern ergänzt sie lediglich. Auch diese Definition ist nicht willkürlich entstanden, sondern so gewählt, dass die bekannten Rechenregeln erhalten bleiben.

Die Aussagen aus Lemma 5.8 gelten auch für alle $a, b \in \mathbb{R}$ und alle $m, n \in \mathbb{Q}$, sofern die Potenzen definiert sind.

Lemma 12.5

Die Gleichungen des folgenden Korollars sind eine direkte Konsequenz aus den Rechenregeln für Potenzen, weil man für $\sqrt[n]{a}$ einfach $a^{1/n}$ schreiben kann. Die Ungleichungen ergeben sich aus Korollar 11.4.

Für alle $a, b \in \mathbb{R}^+ \cup \{0\}$ und alle $m, n \in \mathbb{N}^+$ gelten die folgenden Aussagen:

Korollar 12.6

- (i) $\sqrt[n]{a} \cdot \sqrt[n]{b} = \sqrt[n]{ab}$
- (ii) $\sqrt[n]{a} / \sqrt[n]{b} = \sqrt[n]{a/b}$ für $b \neq 0$
- (iii) $\sqrt[n]{\sqrt[m]{a}} = \sqrt[nm]{a}$
- (iv) $a^{m/n} = (\sqrt[n]{a})^m$
- (v) $a < b \Rightarrow \sqrt[n]{a} < \sqrt[n]{b}$
- (vi) $1 < a \wedge m < n \Rightarrow \sqrt[n]{a} < \sqrt[m]{a}$
- (vii) $a < 1 \wedge m < n \Rightarrow \sqrt[m]{a} < \sqrt[n]{a}$

¹⁴Ein Verfahren, das cleverer als Probieren ist, ist schon seit dem Altertum bekannt: das sogenannte [babylonische Wurzelziehen](#).

Aufgabe 12.18. Berechnen Sie ohne elektronische Hilfsmittel $\sqrt{(-3)^2}$, $\sqrt[3]{27}$, $\sqrt[3]{1/8}$, $\sqrt[4]{32}/\sqrt[8]{4}$, $\sqrt{4/9}$, $\sqrt{2/3} \cdot \sqrt{6}$, $16^{1/4}$, $(\sqrt[3]{5})^2 \cdot \sqrt[6]{25}$ und $25^{-1/2}$. Alle Ergebnisse sind rationale Zahlen. Geben Sie als Antworten ganze Zahlen oder Brüche an, keine Zahlen mit Nachkommastellen.

Aufgabe 12.19. Eine der beiden Aussagen $\forall a \in \mathbb{R} \sqrt{a^2} = a$ und $\forall a \in \mathbb{R}^+ (\sqrt{a})^2 = a$ ist falsch. Welche ist es? Begründen Sie Ihre Antwort mit einem konkreten Gegenbeispiel.

Aufgabe 12.20. Finden Sie für die folgenden Mengen, die alle absichtlich kompliziert aufgeschrieben wurden, jeweils möglichst einfache Darstellungen:

$$\begin{aligned} A &= \{n \in \mathbb{N} : n^2 < 1\} & B &= \{n \in \mathbb{N} : n^2 \geq 0\} \\ C &= \{x \in \mathbb{Q} : x \in \mathbb{R}\} & D &= \{x \in \mathbb{R} : x \in \mathbb{Q}\} \\ E &= \mathbb{N} \cup \mathbb{Q} & F &= \mathbb{N} \setminus \mathbb{Q} \\ G &= \{|z| : z \in \mathbb{Z}\} & H &= \{-z : z \in \mathbb{Z}\} \end{aligned}$$

Aufgabe 12.21. Jeweils zwei der unten aufgeführten Mengen sind gleich. Finden Sie die entsprechenden Paare.

$$\begin{aligned} A &= \mathbb{R} \cap \mathbb{R} & B &= \{x \in \mathbb{Q} : x^2 = 8\} \cup \mathbb{N} \\ C &= \mathbb{R} \cup \mathbb{R} & D &= \{x \in \mathbb{Q} : x^2 = x \wedge x < 1\} \\ E &= \mathbb{N} \cup \{42\} & F &= \{x \in \mathbb{R} : x^2 = x \wedge x < 0\} \\ G &= \{2\sqrt{2}\} & H &= \{x \in \mathbb{R} : x^2 = 8 \wedge x > 0\} \\ I &= \{n \in \mathbb{N} : n^2 \leq 0\} & J &= \{x \in \mathbb{Q} : x^2 = 8 \wedge x > 0\} \end{aligned}$$

Aufgabe 12.22. Jeweils zwei der unten aufgeführten Mengen sind gleich. Finden Sie die entsprechenden Paare.

$$\begin{aligned} A &= \left\{ \frac{m}{n} : m \in \mathbb{N} \wedge n \in \mathbb{N}^+ \right\} & B &= \mathbb{Z} \cap \mathbb{N} \\ C &= \{r \in \mathbb{R} : \exists z \in \mathbb{Z} (z < 0 \wedge r = |z|)\} & D &= \{1, 2, 3, 4, \dots\} \\ E &= \{r \in \mathbb{R} : r < 0 \wedge r^2 = r\} & F &= \mathbb{N} \setminus \{\sqrt{2}, \pi, -42\} \\ G &= \{5y : y \in \mathbb{Q} \wedge y \geq 0\} & H &= \{s \in \mathbb{R} : s < 1 \wedge 1 - s < 0\} \end{aligned}$$

Aufgabe 12.23. Jeweils zwei der unten aufgeführten Mengen sind gleich. Finden Sie die entsprechenden Paare.

$$\begin{aligned} A &= (-2, 4] \cap [-4, 2) & B &= \{x \in \mathbb{R} : x > 0 \vee x < 2\} \\ C &= (-2, 2) \cap \mathbb{Q} & D &= \{x \in \mathbb{R} : x \leq 0 \wedge x \geq 2\} \\ E &= \{x \in \mathbb{Q} : x^2 < 4\} & F &= \{x \in \mathbb{R} : |x| < 4\} \\ G &= (-\infty, 2] \cup [-2, \infty) & H &= \{x \in \mathbb{R} : x^2 < 4\} \\ I &= (-\infty, -2] \cap [2, \infty) & J &= (-4, 4) \end{aligned}$$

Bemerkungen

Beim Vorlesen von periodischen Dezimalzahlen muss man die Periode ansagen, *bevor* sie beginnt. $0.\overline{13}$ ist also „null komma *Periode* eins drei“ und $0.1\overline{3}$ ist „null komma eins *Periode* drei“. Leider hört man des Öfteren so etwas wie „null komma eins drei *Periode*“. Dann ist aber nicht klar, welche von beiden Zahlen gemeint ist.

Die Aussagen aus diesem Abschnitt gelten in entsprechend abgewandelter Form auch dann noch, wenn man statt 10 eine andere Basis wählt. Das ist insbesondere für Computer relevant, die im **Dualsystem** – also mit der Basis 2 – rechnen. Im Dualsystem ergibt sich beispielsweise

$$\frac{1}{10} = 0.0001100_2 = \frac{1}{2^4} + \frac{1}{2^5} + \frac{1}{2^8} + \frac{1}{2^9} + \frac{1}{2^{12}} + \frac{1}{2^{13}} + \dots$$

und es gibt dann noch weniger Zahlen als im Dezimalsystem, die sich mit endlich vielen Nachkommastellen darstellen lassen. (Nämlich nur die Brüche, deren Nenner in gekürzter Form Zweierpotenzen sind.) Eine Konsequenz daraus ist, dass man mit **Fließkommazahlen** (für die nur endlich viele Nachkommastellen zur Verfügung stehen) die meisten reellen Zahlen *nicht* exakt repräsentieren kann.¹⁵

Es sei noch einmal an die Bemerkung am Ende des **Abschnitts über Teilbarkeit** erinnert, dass das Stellenwertsystem zwar sehr praktisch ist, die Darstellung einer Zahl aber nichts an ihren Eigenschaften ändert. 17 ist auch dann noch eine Primzahl, wenn man diese Zahl römisch XVII schreibt. Und $1/5$ ist und bleibt ein Bruch, obwohl diese Zahl im Dezimalsystem nur endlich viele Nachkommastellen, **binär** jedoch unendlich viele hat. (Irrationale Zahlen haben allerdings in jedem Stellenwertsystem unendlich viele Nachkommastellen.)

Es ist empfehlenswert, beim Rechnen mit Wurzeln statt $\sqrt[n]{a}$ immer $a^{1/n}$ zu denken bzw. zu schreiben. Dann kann man mit den bekannten Regeln für Potenzen arbeiten und muss sich keine weiteren Regeln für Wurzeln wie Korollar 12.6 merken.

Es gibt in jedem Semester ein paar Studierende, die berichten, dass man in der Schule $[a, b[$ statt $[a, b)$ geschrieben habe. Wenn Sie das unbedingt so schreiben wollen, dann machen Sie das. (Dann aber konsistent und nicht mal so und mal so!) Ich persönlich finde solche „verdrehten“ Klammern ausgesprochen hässlich und werde sie in diesem Skript nicht verwenden.

Das Zeichen ∞ in Ausdrücken wie $[a, \infty)$ ist nur ein Symbol. ∞ ist *keine* reelle Zahl! Daher würde auch so etwas wie „ $[a, \infty]$ “ keinen Sinn ergeben.¹⁶ Wir **hatten das zwar schon**, aber es schadet sicher nicht, noch einmal darauf hinzuweisen.

Die Intervallschreibweise hat nichts mit der Notation $[n]$ zu tun, die in Abschnitt 4 eingeführt wurde. Bei Intervallen stehen zwischen den Klammern immer *zwei* Zahlen, die durch ein Komma getrennt sind. Eine weitere typische Fehlinterpretation wird in Aufgabe 12.13 angesprochen.

¹⁵Mehr dazu im **nächsten Abschnitt**.

¹⁶Es gibt verschiedene Möglichkeiten, **unendliche Elemente zu $\mathbb{R}$ hinzuzufügen**. Das werden wir in diesem Rahmen jedoch nicht behandeln.

Wir verwenden dieselbe Schreibweise (a, b) für Intervalle und später auch für geordnete Paare. Solche Überschneidungen werden sich nie ganz vermeiden lassen. In der Praxis wird sich die Verwechslungsgefahr als sehr gering herausstellen und zur Not kann man durch eine Formulierung wie „das Intervall (a, b) “ klarstellen, was gemeint ist.

★ Reelle Zahlen und Wurzeln im Computer

Die Quadratwurzel berechnet man in PYTHON mit `sqrt`. Für andere Wurzeln muss man die [bereits erwähnte](#) Potenz `**` verwenden.

```
from math import sqrt
sqrt(49), 27**(1/3)
```

Beachten Sie, dass man die Funktion `sqrt` erst [importieren](#) muss und dass beide Ergebnisse Fließkommazahlen sind, obwohl es sich mathematisch um ganze Zahlen handelt.

Dass Fließkommazahlen intern binär gespeichert werden und deshalb auch Zahlen, die sich dezimal mit wenigen Nachkommastellen exakt darstellen lassen, zu subtilen Fehlern führen können, zeigt ein einfaches Beispiel wie dieses:

```
0.2*3
```

Das ist kein spezifisches Problem von PYTHON. Siehe Fußnote 17.

13. Runden und Fehlerfortpflanzung

In der Praxis rechnet man oft mit ungenauen Werten. Das kann verschiedene Gründe haben:

- Man hat Ausgangsdaten, die durch Messungen oder Schätzungen entstanden sind. Exakte Messungen von physikalischen Größen wie Länge oder Masse sind grundsätzlich unmöglich. Und bei Mengen von Objekten, die zu groß zum Erfassen aller Einzelfälle sind, kann man sich nur mit Schätzungen behelfen.
- Zahlen werden durch die Darstellung und das Rechnen im Computer ungenau. Das liegt daran, dass Computer nur ganz wenige Zahlen exakt darstellen können.¹⁷
- Für viele Berechnung werden [numerische Verfahren](#) eingesetzt, die nur Näherungswerte liefern. Das kann daran liegen, dass man exakte Werte gar nicht angeben kann, aber auch daran, dass die Ermittlung derselben zu aufwendig wäre.

In der Regel gibt es also einen „korrekten“ Wert x und einen Näherungswert, der sich mehr oder weniger von x unterscheiden kann. Die Anführungszeichen stehen

¹⁷Ausführlich wird das in [Konkrete Mathematik \(nicht nur\) für Informatiker](#) erklärt, insbesondere in Kapitel 13 über das Format [IEEE 754](#).

dort, weil x typischerweise unbekannt oder nicht darstellbar ist; sonst müsste man ja nicht mit Näherungen rechnen.

Wenn die Ausgangsdaten und die Berechnung fehlerhaft sind, dann wird auch das Ergebnis nicht exakt sein. Taschenrechner, Computer und andere technische Geräte geben jedoch eine konkrete Zahl aus, die man in so einer Situation nicht kritiklos weitergeben darf. Sinnvoll wäre im Allgemeinen die Angabe von Fehlerschranken. Wir werden uns allerdings nicht ausführlich mit der sogenannten **Fehlerrechnung** beschäftigen. Es geht zunächst nur darum, ein Bewusstsein für „zu genaue Zahlen“ zu entwickeln.

Nehmen wir als konkretes Beispiel an, dass Sie für eine handwerkliche Arbeit die Fläche eines Kreises berechnen müssen. Sie messen als Radius 84 Zentimeter aus, geben in PYTHON

```
from math import pi

r = 0.84
pi*r**2
```

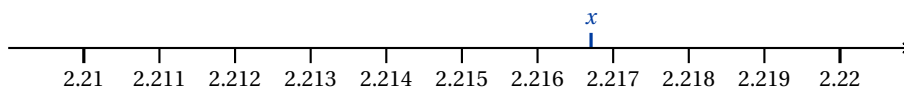
ein und erhalten als Ergebnis: 2.216707776372958. Aber was heißt denn das? Die beiden Angaben 2.216707776372 und 2.216707776373 unterscheiden sich beispielsweise nur um ein Quadratmikrometer. Ihre Messung stimmte aber höchstwahrscheinlich nur bis auf einem Zentimeter oder maximal einen Millimeter. Daher ergibt es keinen Sinn, so viele Nachkommastellen anzugeben. Man *rundet* das Ergebnis.

Wenn eine Zahl $x \in \mathbb{R}$ auf n Stellen nach dem Komma **gerundet** wird, dann ist damit gemeint, dass man eine Zahl $\tilde{x}$ angibt, die zwei Bedingungen erfüllt:

- (i) In der Dezimaldarstellung von $\tilde{x}$ sind höchstens die ersten n Nachkommastellen von null verschieden.
- (ii) Der Betrag $|\tilde{x} - x|$ der Abweichung vom richtigen Wert ist nicht größer als $5 \cdot 10^{-(n+1)}$.

Definition

Ist x die Zahl 2.216707776372958 von oben und will man auf drei Stellen nach dem Komma runden, so soll die Abweichung nicht größer als $5 \cdot 10^{-4} = 0.0005$ werden. Das ist die Hälfte des Abstands zweier Zahlen in der folgenden Skizze.



Der gerundete Wert kann also nur 2.217 sein. Offensichtlich ist durch die obige Definition der gerundete Wert nur dann nicht eindeutig bestimmt, wenn x genau zwischen zwei „Kandidaten“ liegt. Will man beispielsweise $x = 0.265$ auf zwei Stellen nach dem Komma runden, so sind sowohl 0.26 als auch 0.27 in Ordnung. Es gibt dafür jedoch zwei Konventionen:

kaufmännisches
Runden

- Beim *kaufmännischen Runden* wird „nach oben“ gerundet. Im obigen Beispiel würde man sich also für 0.27 entscheiden.¹⁸

mathematisches
Runden

- Beim *wissenschaftlichen* oder *mathematischen Runden* wird so gerundet, dass die letzte Ziffer (nach dem Runden) gerade ist. Dann käme oben 0.26 heraus.

Man beachte, dass beide Konventionen nur dann greifen, wenn es eine Wahlmöglichkeit gibt! Im ersten Beispiel kommt *immer* 2.217 heraus, egal ob man kaufmännisch oder mathematisch rundet.

Der Sinn des mathematischen Rundens, das Computer typischerweise auch intern einsetzen, ist das Vermeiden systematischer Abweichungen. Man kann sich das anhand einer Summe wie

$$0.15 + 0.25 + \dots + 0.85 + 0.95$$

vor Augen führen. Die korrekte Summe ist 4.95. Würde man die neun Summanden vor der Addition kaufmännisch auf eine Stelle nach dem Komma runden, so käme man auf 5.4 und läge um 0.45 daneben, weil immer in dieselbe Richtung gerundet wurde. Mit mathematischem Runden kommt man auf den wesentlich besseren Wert 5.0, weil sich die Rundungsfehler größtenteils gegenseitig neutralisieren.

Aufgabe 13.1. Geben Sie den Wert von π gerundet auf drei, vier und sieben Stellen nach dem Komma an. Geben Sie den Wert von $3/80$ gerundet auf drei Stellen nach dem Komma an.

Aufgabe 13.2. Warum wird man sich beim Runden von $\sqrt{2}$ niemals die Frage stellen müssen, ob man kaufmännisch oder mathematisch runden sollte?

Aufgabe 13.3. In Aufgabe 13.1 hat man gesehen, dass beispielsweise beim Runden auf drei Nachkommastellen die dritte Nachkommastelle des gerundeten Wertes nicht die dritte Nachkommastelle von π war. Geben Sie ein Beispiel für eine Zahl an, bei der nach Runden auf drei Nachkommastellen *keine* Nachkommastelle mehr mit dem Original übereinstimmt.

Aufgabe 13.4. Falls Sie noch nie so richtig über Runden und Näherungswerte Gedanken gemacht haben, sollten Sie sich vielleicht die Videos *Genau oder ungefähr?* und *Wenn Zahlen zu genau sind ...* anschauen.

Die Angabe von Nachkommastellen ist oft nicht zielführend. Bei physikalischen Größen hängt deren Relevanz z. B. von der Maßeinheit ab. Eine Nachkommastelle bei einer Angabe in Tonnen ist etwas ganz anderes als eine Nachkommastelle bei einer Angabe in Gramm. Daher spricht man auch von *signifikanten Stellen*, die unabhängig von solchen Skalierungen sind.

¹⁸Diese Regel wird im Allgemeinen nur für positive Zahlen angewendet.

Wenn eine Zahl $x \in \mathbb{R}_{>0}$ auf n **signifikante Stellen** nach dem Komma gerundet wird, dann ist damit gemeint, dass man eine Zahl $\tilde{x}$ angibt, die folgendermaßen konstruiert wird:

- (i) Es gibt ein eindeutig bestimmtes $k \in \mathbb{Z}$ mit $0 \leq 10^k x < 1 \leq 10^{k+1} x$.
- (ii) $\tilde{x}$ wird so gewählt, dass $10^k \tilde{x}$ der auf n Nachkommastellen gerundete Wert von $10^k x$ ist.

Für negative x geht man entsprechend vor. Die Konventionen für kaufmännisches und mathematisches Runden sind dieselben wie oben.

Definition

Im ersten Schritt wird die erste Ziffer von x , die von null verschieden ist, direkt hinter das Dezimaltrennzeichen verschoben. Aus $x_1 = 0.000346$ wird mit $k = 3$ z. B. 0.346 und aus $x_2 = 42.81$ wird mit $k = -2$ der Wert 0.4281. Dann wird auf Nachkommastellen gerundet und anschließend die Verschiebung rückgängig gemacht. x_1 auf zwei signifikante Stellen gerundet ist 0.00035, x_2 auf zwei signifikante Stellen gerundet ist 43.

In der Regel wird so gerundet, dass man den verbleibenden Stellen „trauen“ kann. Daher ist es üblich, „überflüssige“ Nullen auszuschreiben, wenn sie beim Runden entstanden sind.

Damit ist gemeint, dass $x = 0.198$ auf zwei Nachkommastellen gerundet 0.20 ist und *nicht* 0.2. Und $\sqrt{24}$ auf drei signifikante Stellen gerundet ist 4.90 und *nicht* 4.9. Mathematisch macht das zwar keinen Unterschied, aber die Signalwirkung beim Lesen ist eine andere. Außerdem sollte man, wenn es sich um Näherungswerte handelt, ein anderes Zeichen als $=$ verwenden, weil das (zumindest in der Mathematik) für exakte Gleichheit steht. Üblich ist in so einem Fall das Symbol $\approx$.

 $\approx$

Ist $\tilde{x}$ ein Näherungswert für die Zahl $x \in \mathbb{R}$, so bezeichnet man die Differenz $\tilde{x} - x$ als den **absoluten Fehler** von $\tilde{x}$.

Ist $\tilde{x}$ ein Näherungswert für die Zahl $x \in \mathbb{R} \setminus \{0\}$, so bezeichnet man den Quotienten $(\tilde{x} - x)/x$ als den **relativen Fehler** von $\tilde{x}$.

Definition

Runden auf n Nachkommastellen stellt also sicher, dass der Betrag des absoluten Fehlers höchstens $5 \cdot 10^{-n-1}$ ist. Runden auf signifikante Stellen sorgt hingegen dafür, dass der relative Fehler unabhängig von der Wahl der Maßeinheit immer gleich ist. Da es sich bei relativen Fehlern um Verhältnisse handelt, werden diese oft in Prozent angegeben.

Zum Thema **Fehlerfortpflanzung** nur zwei einfache Beispiele, um deutlich zu machen, dass auch bei fehlerfreier Rechnung einmal vorhandene Abweichungen immer größer werden können.

Seien x und y reelle Zahlen und $\tilde{x}$ und $\tilde{y}$ Näherungswerte für diese Zahlen mit absoluten Fehlern Δ_x bzw. Δ_y und relativen Fehlern δ_x bzw. δ_y .

Lemma 13.1

Betrachtet man $\tilde{x} \pm \tilde{y}$ als Näherung für $x \pm y$, so ist der Betrag des absoluten Fehlers höchstens $|\Delta_x| + |\Delta_y|$. Betrachtet man $\tilde{x} \cdot \tilde{y}$ als Näherung für $x \cdot y$, so ist der Betrag des relativen Fehlers höchstens $|\delta_x| + |\delta_y| + |\delta_x||\delta_y|$. Die angegebenen Schranken lassen sich nicht verbessern, d. h., im schlimmsten Fall können die Fehler tatsächlich so groß wie angegeben sein.

Aufgabe 13.5. Geben Sie Beispiele dafür an, wie in Lemma 13.1 die Fehler maximal groß werden.

Bemerkungen

In der (reinen) Mathematik kommt im Gegensatz zu den Natur- und Ingenieurwissenschaften die Schreibweise von Zahlen mit Nachkommastellen so gut wie nie vor. Sie sollten grundsätzlich immer Brüche verwenden, wenn dies möglich ist. Schreiben Sie also beispielsweise $3/4$ und *nicht* 0.75, weil Letzteres einen ungenauen Wert suggeriert.

Beachten Sie, dass irrationale Zahlen ohnehin in dieser Schreibweise nicht korrekt geschrieben werden können. 0.75 statt $3/4$ ist nur schlechter Stil, 1.73205 statt $\sqrt{3}$ ist jedoch schlicht und einfach falsch!

★ Typische Fehlerfortpflanzung im Computer

Im Folgenden will ich an zwei Beispielen Probleme mit der Arithmetik im Computer darstellen. Die Berechnungen werden zwar in PYTHON vorgeführt, in anderen Programmiersprachen kommt es jedoch zu denselben Ergebnissen, da Fließkommaarithmetik von der CPU durchgeführt wird und *standardisiert* ist.

```
from math import sqrt

x = sqrt(2)
m = 1000000
s, S = 0, 0
for k in range(m):
    s += x
for k in range(m):
    S += s
S*x - 2*m*m
```

Hier wurde $x \cdot \sum_{k=1}^m (\sum_{k=1}^m x)$ berechnet. Das ist dasselbe wie $m^2 x^2$. Der Unterschied zwischen beiden Werten im Computer ist jedoch größer als 45.

```
s0, s1 = 0, 0
for k in range(0, 14):
    s0 += 9*10**(-k)
    s1 += 9*10**(k-13)
s0, s1
```

Und hier wurde $\sum_{k=0}^{13} 9/10^{-k}$ auf zwei verschiedene Arten (einmal „vorwärts“ und einmal „rückwärts“) berechnet. Es kommen zwei unterschiedliche Ergebnisse heraus. (Welches ist richtig?)

Die Beispiele kommen Ihnen vielleicht konstruiert vor. Das sind sie jedoch nur, um die Probleme besonders deutlich hervorzuheben. Grundsätzlich sollten Sie davon ausgehen, dass mit dem Computer berechnete Fließkommawerte *immer* falsch sind. Exakt richtig sind sie nur in Ausnahmefällen. Daher ist es auch keine gute Idee, solche Zahlen auf Gleichheit zu testen. Probieren Sie z. B. einmal aus, welche Antwort Sie auf die Eingabe

```
6*0.1 == 0.6
```

erhalten. Für vertiefende Informationen zu diesem Thema siehe Fußnote 17.

14. Die komplexen Zahlen

Komplexe Zahlen haben sich im 16. Jahrhundert gleichsam in die Mathematik „eingeschlichen“. ¹⁹ Inzwischen sind sie aber ein unverzichtbarer Bestandteil nicht nur der Mathematik, sondern auch der Naturwissenschaften und der Informatik. So kann man z. B. die Programmierung von [Quantencomputern](#) ohne komplexe Zahlen nicht verstehen. ²⁰ Wir werden bei der Einführung der komplexen Zahlen in gewissem Sinne der historischen Entwicklung (allerdings in moderner Schreibweise) folgen, haben dabei jedoch den Vorteil, dass wir bereits wissen, dass es ein „Happy End“ gibt.

Wegen Punkt (vii) von Lemma 11.3 ist klar, dass Gleichungen wie $x^2 = -1$ keine reelle Lösung haben können. Wir fügen daher neue Zahlen hinzu, die die gewünschte Eigenschaft haben.

Ein Objekt der Form $z = a + bi$ mit reellen Zahlen a und b nennt man eine **komplexe Zahl** mit **Realteil** $\operatorname{Re}(z) = a$ und **Imaginärteil** $\operatorname{Im}(z) = b$. Für die Menge $\{a + bi : a, b \in \mathbb{R}\}$ aller komplexen Zahlen schreiben wir $\mathbb{C}$. Das Symbol i wird **imaginäre Einheit** genannt.

Definition

Damit sind $3 + 5i$, $\frac{1}{3} + (-8)i$ oder $-\pi + \sqrt{2}i$ Beispiele für komplexe Zahlen. Für die zweite Zahl werden wir jedoch in Zukunft kürzer $\frac{1}{3} - 8i$ schreiben. Der Realteil dieser Zahl ist $\frac{1}{3}$ und ihr Imaginärteil ist -8 . Das von [Carl Friedrich Gauß](#), eingeführte Adjektiv *komplex* wird hier im Sinne von „zusammengesetzt“ (und nicht etwa „kompliziert“) verwendet, weil solche Zahlen immer zwei „Teile“ haben.

Zahlen ergeben jedoch nur einen Sinn, wenn man mit ihnen rechnen kann. Wir legen nun fest, wie das geschehen soll. Wenn man sicher im Umgang mit der [Arithmetik](#) ist, muss man eigentlich gar nichts Neues lernen.

¹⁹Siehe dazu mein Video [Woher die komplexen Zahlen kommen](#).

²⁰Zu Quantencomputern gibt es [eine sechsteilige Videoreihe](#) von mir.

Definition

Komplexe Zahlen der Form $a + bi$ werden beim Rechnen mit den vier Grundrechenarten so behandelt, als wäre i eine reelle Zahl wie a und b . Der einzige Unterschied ist, dass $i^2 = -1$ gilt.

Wir gehen das an Beispielen durch. Addition und Subtraktion sind besonders einfach (wenn man mit den Vorzeichen nicht durcheinanderkommt).

$$(3 + 2i) + (5 - 4i) = 3 + 2i + 5 - 4i = 3 + 5 + 2i - 4i = 8 - 2i$$

$$(3 + 2i) - (5 - 4i) = 3 + 2i - 5 + 4i = 3 - 5 + 2i + 4i = -2 + 6i$$

Bei der Multiplikation spielt zum ersten Mal die besondere Eigenschaft der imaginären Einheit eine Rolle.

$$\begin{aligned}(3 + 2i) \cdot (5 - 4i) &= 3 \cdot 5 + (2i) \cdot 5 + 3 \cdot (-4i) + (2i) \cdot (-4i) \\ &= 15 + 10i + (-12i - 8i^2) \stackrel{!}{=} 15 + 10i + (-12i - 8 \cdot (-1)) \\ &= 15 + 10i + (-12i + 8) = 23 - 2i\end{aligned}$$

Beachten Sie, dass die Klammern am Anfang (wegen der Regel „Punkt- vor Strichrechnung“) wirklich nötig sind, während sie bei den beiden Beispielen davor nur aus didaktischen Gründen vorhanden waren.

Aufgabe 14.1. In den Beispielen verwenden wir der Übersichtlichkeit halber oft ganzzahlige Real- und Imaginärteile. Man darf jedoch nicht vergessen, dass alle reellen Zahlen erlaubt sind. Berechnen Sie $z + w$, $z - w$ und zw für $z = \frac{2}{3} + \frac{5}{8} \cdot i$ und $w = \frac{3}{7} - \frac{5}{2} \cdot i$. Berechnen Sie außerdem die Summe von $\sqrt{2} - 3i$ und $5 - \pi i$.

Aufgabe 14.2. Geben Sie Real- und Imaginärteil von $z = 3 - \sqrt{5}i$ an.

Das Ziel bei solchen Rechnungen ist, das Ergebnis so darzustellen, dass Real- und Imaginärteil klar erkennbar sind. Bei der Division ist dafür ein kleiner Trick vonnöten, für den wir zunächst ein Konzept einführen, das wir noch brauchen werden.

Definition

Ist $z = a + bi$ eine komplexe Zahl, so bezeichnen wir die Zahl $\bar{z} = a - bi$ als die zu z **konjugiert** komplexe Zahl.

$\overline{3 + 4i}$ ist also $3 - 4i$ und $\overline{3 - 4i}$ ist $3 + 4i$. Insbesondere gilt offenbar immer $\overline{\bar{z}} = z$. Sie werden zugeben müssen, dass das simpel ist. Es ist allerdings auch sehr hilfreich, denn für jede komplexe Zahl $z = a + bi$ gilt

$$z \cdot \bar{z} = (a + bi)(a - bi) = a^2 - b^2 i^2 = a^2 + b^2 \quad (14.1)$$

und damit ist $z \cdot \bar{z}$ immer eine nichtnegative reelle Zahl.

Der besagte Trick ist nun, bei der Division zweier komplexer Zahlen mit der komplex konjugierten Zahl des Divisors zu erweitern, weil dann danach „unten“ nur noch eine reelle Zahl steht.

$$\frac{3+2i}{5-4i} = \frac{3+2i}{5-4i} \cdot \frac{5+4i}{5+4i} = \frac{7+22i}{41} = \frac{7}{41} + \frac{22}{41} \cdot i$$

Dadurch erhalten wir wieder eine Darstellung, in der Real- und Imaginarteil sauber voneinander getrennt sind. In Zukunft werden wir übrigens solche Zahlen eher in der Form $\frac{7}{41} + \frac{22i}{41}$ oder $\frac{7+22i}{41}$ bzw. $(7+22i)/41$ schreiben.

Aufgabe 14.3. Berechnen Sie z/w und w/z für die Zahlen aus Aufgabe 14.1.

Aufgabe 14.4. Wenn die komplexen Zahlen neu für Sie sind, dann sollten Sie das Rechnen mit ihnen auf jeden Fall üben. In [Wolfram Alpha](#) können Sie Ausdrücke wie $(4+i)/(2-5i)$ direkt eingeben und damit Ihre Ergebnisse überprüfen.

Aufgabe 14.5. Sei $w = a + bi$ und $z = c + di$. Berechnen Sie $\overline{w+z}$ und $\overline{w} + \overline{z}$. Was fällt Ihnen auf? Und wie ist es mit $\overline{w \cdot z}$ und $\overline{w} \cdot \overline{z}$?

Aufgabe 14.6. Was kommt heraus, wenn man $z + \overline{z}$ bzw. $z - \overline{z}$ berechnet?

Was wir an den Beispielen herausgearbeitet haben, fassen wir noch einmal zusammen:

Summen, Differenzen, Produkte und Quotienten von komplexen Zahlen sind wieder komplexe Zahlen. $\mathbb{R}$ ist eine Teilmenge von $\mathbb{C}$. Zwei komplexe Zahlen $a_1 + b_1i$ und $a_2 + b_2i$ sind genau dann gleich, wenn $a_1 = a_2 \wedge b_1 = b_2$ gilt.

Lemma 14.1

Insbesondere sind also alle reellen Zahlen komplexe Zahlen, denn eine reelle Zahl a kann man als komplexe Zahl $a + 0i$ mit dem Imaginärteil 0 darstellen. Komplexe Zahlen, die *nicht* reell sind (deren Imaginärteil also nicht verschwindet), nennt man auch *echt komplex*.²¹ Komplexe Zahlen der Form $0 + bi$, deren Realteil verschwindet, werden *rein imaginär* genannt. (Die Zahl null ist also die einzige Zahl, die sowohl reell als auch rein imaginär ist.)

echt komplex
rein imaginär

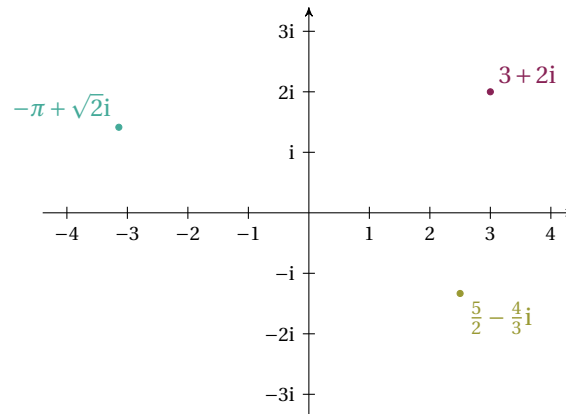
Der letzte Teil von Lemma 14.1 ist nicht so selbstverständlich, wie man vielleicht denken könnte, denn für Brüche gibt es ja unendlich viele verschiedene Darstellungen: Aus $a_1/b_1 = a_2/b_2$ folgt eben *nicht* $a_1 = a_2 \wedge b_1 = b_2$, d. h., Zähler und Nenner sind nicht eindeutig bestimmt.²² Real- und Imaginärteil einer komplexen Zahl sind jedoch eindeutig bestimmt.

Bevor wir dazu kommen, wozu die komplexen Zahlen eigentlich gut sind, zunächst eine schlechte Nachricht: Es ergibt keinen Sinn, sich die komplexen Zahlen analog zu den reellen auf so etwas wie einer „komplexen Zahlengeraden“ vorzustellen.

²¹Analogie: Alle rationalen Zahlen sind reelle Zahlen. Reelle Zahlen, die *nicht* rational sind, nennt man *irrational*.

²²Es sei denn, man geht explizit von der gekürzten Darstellung aus und legt fest, dass beispielsweise der Nenner nie negativ sein soll.

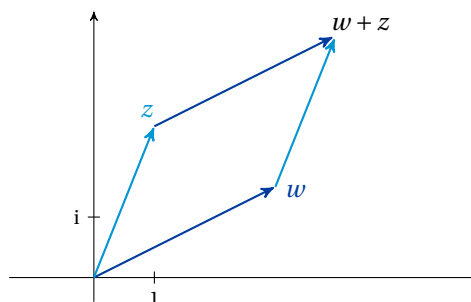
Man könnte die komplexen Zahlen zwar irgendwie „sortieren“, aber man kann beweisen, dass keine solche Sortierung sich im Sinne von Abschnitt 11 mit den Rechenoperationen vertragen würde.²³ Komplexe Zahlen können deswegen auch nicht „positiv“ oder „negativ“ sein.²⁴ Es gibt jedoch eine geometrische Interpretation der komplexen Zahlen, die sehr hilfreich ist: Man trägt die Zahlen $a + bi$ als Punkte (a, b) in ein **kartesisches Koordinatensystem** ein.



gaußsche Zahlenebene

Die reelle Zahlengerade bildet die horizontale Achse und der Rest der Ebene ist mit echt komplexen Zahlen bedeckt. Auf der vertikalen Achse liegen die rein imaginären Zahlen. Man nennt das die **gaußsche Zahlenebene**.

Das ist nicht nur eine hübsche Darstellung, sondern sie liefert auch sinnvolle geometrische Gegenstücke zu den Rechenoperationen. Für die Addition zweier komplexer Zahlen w und z stellt man sich beide als Pfeile vor, die von 0 (also vom Ursprung des Koordinatensystems) auf diese Zahlen zeigen. Verschiebt man den z -Pfeil so, dass sein Anfang auf der Spitze des w -Pfeils landet, so zeigt seine Spitze auf $w + z$.²⁵ Das kennen Sie aus der Schule vom Rechnen mit **Vektoren** bzw. vom **Kräfteparallelogramm**.



Aufgabe 14.7. Wenn Sie die Lage einer komplexen Zahl z in der gaußschen Zahlenebene kennen, wo liegt dann $-z$?

²³Eine Begründung dafür findet man im Abschnitt über **geordnete Körper** im **alten Skript**.

²⁴Das hat u. a. die Konsequenz, dass Sie, wenn Sie in einem Mathebuch so etwas wie $x > 0$ lesen, sich sicher sein können, dass x *nicht* komplex ist.

²⁵Und mit vertauschten Rollen erhält man dasselbe Ergebnis, wodurch man gleich eine visuelle Bestätigung der Kommutativität der Addition erhält.

Aufgabe 14.8. Wenn Sie die Lage einer komplexen Zahl z in der gaußschen Zahlenebene kennen, wo liegt dann $\bar{z}$?

Aufgabe 14.9. Zeichnen Sie eine ähnliche Skizze wie oben für die Subtraktion zweier komplexer Zahlen.

Wie gesagt sollten Ihnen die graphische Addition und Subtraktion komplexer Zahlen (wenn auch nicht unter diesem Namen) vertraut sein. Aber wie kann man sich die Multiplikation vorstellen? Dafür müssen wir noch ein wenig Vorarbeit leisten, die wir ohnehin brauchen werden. Wie berechnet man den Abstand eines Punktes (a, b) vom Ursprung? Genau, mit dem **Satz des Pythagoras**. Der Abstand ist die Quadratwurzel von $a^2 + b^2$. Dieser Zahl geben wir einen Namen.

Der **Absolutbetrag** bzw. **Betrag** einer komplexen Zahl $z = a + bi$ ist die Zahl $|z| = \sqrt{a^2 + b^2} = \sqrt{\operatorname{Re}(z)^2 + \operatorname{Im}(z)^2}$.

Definition

Man beachte, das z zwar irgendeine komplexe Zahl sein kann, $a^2 + b^2$ aber immer eine nichtnegative reelle Zahl ist. Die Quadratwurzel ist also immer definiert und ebenfalls eine nichtnegative reelle Zahl (wie es sich für einen Abstand gehört.)

Aufgabe 14.10. Wo kam kürzlich der Ausdruck $a^2 + b^2$ schon einmal vor?

Aufgabe 14.11. Berechnen Sie $|i|$.

Aufgabe 14.12. Sei $z = a + bi$. Vergleichen Sie $|z|$ und $|-z|$.

Aufgabe 14.13. Welcher Zusammenhang besteht zwischen dem komplexen Betrag und dem auf Seite 44 definierten reellen Betrag?

Aufgabe 14.14. Sei $w = a + bi$ und $z = c + di$. Berechnen Sie $|wz|$ und $|w| \cdot |z|$. Was fällt Ihnen auf?

Der komplexe Betrag hat die Eigenschaften, die für den reellen Betrag schon in Lemma 5.7 genannt wurden:

Für alle $w, z \in \mathbb{C}$ gilt:

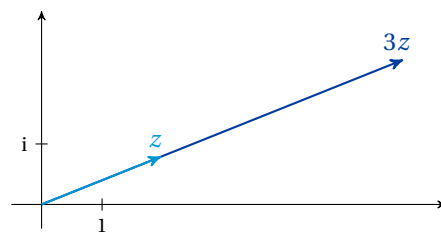
- (i) $|wz| = |w| \cdot |z|$
- (ii) $|w/z| = |w|/|z|$ für $z \neq 0$
- (iii) $|w + z| \leq |w| + |z|$
- (iv) $|w - z| \leq |w| + |z|$
- (v) $||w| - |z|| \leq |w \pm z|$

Lemma 14.2

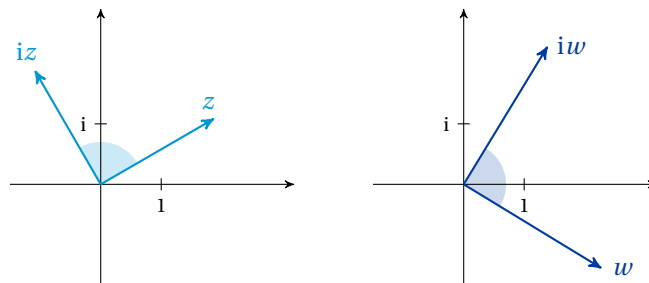
In Aufgabe 14.14 haben Sie sich selbst bereits von der ersten Eigenschaft überzeugt. Auch der Rest lässt sich leicht überprüfen. Lediglich für Eigenschaft (iii), die wir *Dreiecksungleichung* genannt hatten, ist etwas mehr Rechenarbeit nötig, die wir uns jedoch ersparen, weil man sie geometrisch leicht einsieht – siehe Aufgabe 14.15.

Aufgabe 14.15. In **Wolfram Alpha** können Sie mit Eingaben wie $|3+5i|$ Beträge von komplexen Zahlen ausrechnen. Berechnen Sie $|w| + |z|$ und $|w + z|$ für $w = 3 + i$ und $z = 1 + 2i$, um sich zu überzeugen, dass man in der Dreiecksungleichung $\leq$ nicht durch $=$ ersetzen kann. (Siehe auch Aufgabe 5.17.) Zeichnen Sie eine Skizze, in der man w , z und $w + z$ sieht, und erklären Sie daran den Namen *Dreiecksungleichung*.

Nun kommen wir zur Multiplikation und schauen uns zunächst zwei besonders einfache Fälle an, von denen der erste auch eine Entsprechung beim Rechnen mit Vektoren hat. Multipliziert man nämlich eine komplexe Zahl $z = a + bi$ mit einer positiven *reellen* Zahl r , so erhält man $rz = ra + rbi$. Das entspricht offenbar einfach einer Streckung ($r > 1$) bzw. Stauchung ($r < 1$) des entsprechenden Pfeils um diesen Faktor. Bei Vektoren nennt man das skalare Multiplikation. Insbesondere hat rz den Abstand $|rz| = r \cdot |z|$ von null, ist also r -mal so „lang“ wie z .²⁶



Multipliziert man $z = a + bi$ mit i , so ergibt sich $iz = -b + ai$. Die Komponenten des zugehörigen Punktes in der gaußschen Zahlenebene werden also vertauscht und bei der hinteren ändert sich zudem bei diesem Tausch das Vorzeichen. Geometrisch ist das eine Drehung um einen rechten Winkel (wobei sich der Abstand von der Null nicht ändert, weil i den Betrag eins hat).

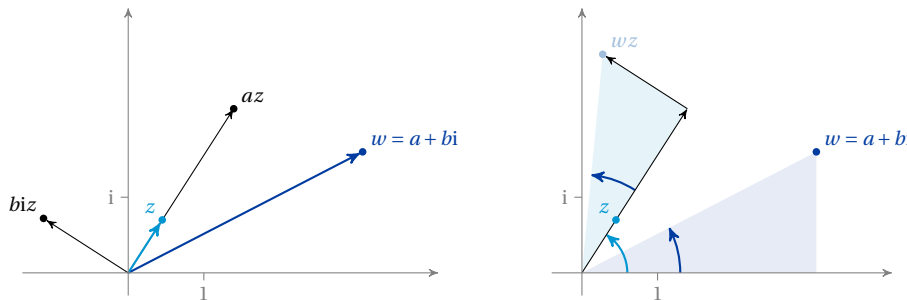


Jetzt schauen wir uns die Multiplikation einer beliebigen komplexen Zahl z mit irgendeiner anderen komplexen Zahl $w = a + bi$ an. (Der Einfachheit halber gehen wir von positiven Zahlen a und b aus. Den Fall, dass a oder b oder beide negativ sind, sollten Sie danach zur Kontrolle selbst skizzieren.) Wir zerlegen das Produkt als $wz = (a + bi)z = az + biz$. Wir haben uns vorher überlegt:

- (i) az entspricht einer Streckung von z um den Faktor a .
- (ii) iz entspricht einer Drehung von z um einen rechten Winkel.
- (iii) biz entspricht einer Streckung von iz um den Faktor b .
- (iv) Die Resultate aus dem ersten und dritten Schritt werden addiert.

²⁶Ist r negativ, so ändert sich außerdem noch die Richtung und rz ist $|r|$ -mal so „lang“ wie z .

Das sieht im Ergebnis so aus:



Mit Schulkenntnissen aus der **Elementargeometrie** können wir nun die folgenden Beobachtungen machen:

- Wir sehen zwei **rechtwinklige Dreiecke**.
- Das rechte hat die Kathetenlängen a und b .
- Das linke hat die Kathetenlängen $a|z|$ und $b|z|$.
- Also sind die beiden Dreiecke **ähnlich** und das linke ist aus dem rechten durch Streckung (oder Stauchung) um den Faktor $|z|$ hervorgegangen.
- Darum sind erstens alle Winkel in den Dreiecken gleich und zweitens hat die Hypotenuse des linken Dreiecks die Länge $|z||w|$ (weil die Hypotenuse des rechten Dreiecks die Länge $|w|$ hat).

Dass $|w||z|$ der Betrag des Produktes wz ist, wussten wir bereits, aber es ist beruhigend, das noch einmal bestätigt zu bekommen. Neu ist jedoch die sehr wichtige Erkenntnis über die Winkel, für die wir einen Begriff einführen.

Für $z \in \mathbb{C} \setminus \{0\}$ sei das **Argument** bzw. die **Phase** von z , geschrieben $\arg(z)$, der **Winkel** zwischen der positiven reellen Achse und der Verbindungslinie von z mit dem Nullpunkt. Um eine eindeutige Angabe zu garantieren, legen wir fest, dass das Argument (in **Bogenmaß**) immer ein Winkel aus dem Intervall $(-\pi, \pi]$ sein soll.

Definition

Damit können wir festhalten: *Das Argument von wz ist die Summe der Argumente von w und z .* Wir erhalten dadurch übrigens auch eine schöne geometrische Begründung für die Regel „minus mal minus ist plus“, die wir schon aus der Schule kennen: Negative reelle Zahlen haben das Argument π (also 180 Grad). Multipliziert man zwei solche Zahlen, so hat das Ergebnis das Argument 2π bzw. 360 Grad, was bedeutet, dass es auf der positiven Hälfte der reellen Achse liegen muss.

Aufgabe 14.16. Sei $z = -a + bi$ mit $a, b > 0$. In welcher Hälfte der gaußschen Zahlenebene liegt z^2 mit Sicherheit *nicht*?

Aufgabe 14.17. Sei $z \in \mathbb{C}$ mit $\operatorname{Re}(z) > 0$ und $\operatorname{Im}(z) < 0$. Was kann man über die Lage von z^3 in der gaußschen Zahlenebene aussagen?

Das Argument von w^2 erhält man also, indem man das Argument von w verdoppelt. Daraus folgt, wie man eine Zahl w mit $w^2 = z$ findet: Der Betrag von w muss $\sqrt{|z|}$ sein und das Argument von w muss die Hälfte des Arguments von z betragen. Zumindest geometrisch ist also klar, dass es so eine „Wurzel“ immer geben muss. (Und damit für $z \neq 0$ sogar zwei, weil $(-w) \cdot (-w) = w \cdot w$ gilt.)

Satz 14.3

Für jede komplexe Zahl $z \in \mathbb{C} \setminus \{0\}$ hat die Gleichung $w^2 = z$ genau zwei verschiedene Lösungen, wobei die zweite Lösung sich durch $-w_1$ ergibt, wenn w_1 die erste ist.

Wenn man weiß, dass solche Gleichungen immer lösbar sind, hat man auch einen Ansatz, mit dem man garantiert zu einem Ergebnis kommt. Das gehen wir an einem Beispiel durch. Eine Lösung der Gleichung $w^2 = 3 + 4i$ muss die Form $w = p + qi$ mit reellen (!) Zahlen p und q haben. Mit $w^2 = (p^2 - q^2) + 2pqi$ erhält man damit die zwei Gleichungen

$$p^2 - q^2 = 3 \quad \text{und} \quad (14.2)$$

$$2pq = 4. \quad (14.3)$$

(14.3) formt man zu $q = 2/p$ um und setzt das in (14.2) ein. Das liefert $p^2 - 4/p^2 = 3$ und nach Multiplikation mit p^2 wird daraus

$$p^4 - 3p^2 - 4 = 0. \quad (14.4)$$

Nun nimmt man die Substitution $p^2 = r$ vor und aus (14.4) wird $r^2 - 3r - 4 = 0$. Das ist eine aus der Schule bekannte **quadratische Gleichung** mit den beiden Lösungen $r_1 = -1$ und $r_2 = 4$. Wegen $p^2 = r$ kommt jedoch nur die zweite Lösung infrage²⁷ und Rücksubstitution liefert für p die Möglichkeiten $p_1 = 2$ und $p_2 = -2$. Mit $q = 2/p$ erhält man $q_1 = 1$ und $q_2 = -1$. Das liefert die Lösungen $w_1 = 2 + i$ und $w_2 = -2 - i$. (Es hätte gereicht, w_1 zu berechnen, weil sich w_2 als $-w_1$ ergeben muss.)

Beachten Sie, dass man bei solchen Aufgaben ganz leicht eine Probe machen kann: Wenn Sie $(2 + i)^2$ berechnen, muss sich natürlich $3 + 4i$ ergeben.

Aufgabe 14.18. Im Beispiel eben wurden zwei Umformungen vorgenommen, bei denen man aufpassen muss. Welche waren das?

Aufgabe 14.19. **Wolfram Alpha** kann natürlich auch Gleichungen wie $w^2 = -7 + 24i$ lösen. Wenn Sie aber einfach irgendwelche komplexen Zahlen einsetzen, werden die Ergebnisse selten „schön“ sein. Zum Üben würde ich das folgende Vorgehen vorschlagen: Berechnen Sie für ein paar komplexe Zahlen z_1 bis z_n die Quadrate z_1^2 bis z_n^2 und schreiben Sie die beiden Listen auf zwei verschiedene Zettel. Am nächsten Tag haben Sie die Zahlen sicher vergessen. Sie können dann aus der zweiten Liste eine Zahl z_k^2 entnehmen, die Aufgabe $w^2 = z_k^2$ lösen und danach Ihr Ergebnis w mit dem Wert z_k aus der ersten Liste vergleichen. Arbeiten Sie ruhig auch ab und zu mit Brüchen und nicht nur mit ganzen Zahlen!

²⁷Weil p reell sein soll!

Ist es nun also OK, $2 + i = \sqrt{3 + 4i}$ zu schreiben? Leider ist die Sache etwas kompliziert. Wenn a eine positive reelle Zahl ist, dann hat die Gleichung $x^2 = a$ zwei Lösungen, von denen eine positiv und eine negativ ist. Mit $\sqrt{a}$ ist dann *per definitionem* immer die positive der beiden Lösungen gemeint! Wenn die komplexen Zahlen ins Spiel kommen, kann man aber nicht mehr von positiv oder negativ reden. Man könnte nun willkürlich etwas festlegen (z. B. die Lösung, die in der rechten Hälfte der gaußschen Zahlenebene liegt), aber keine solche Festlegung liefert eine *Funktion*, die für *alle* komplexen Zahlen definiert ist und „angenehme“ mathematische Eigenschaften wie *Stetigkeit* oder *Differenzierbarkeit* hat. Auf weitere Details kann ich hier nicht eingehen und belasse es bei dem Hinweis, dass Sie *nicht* $\sqrt{z}$ schreiben sollten, wenn z eine beliebige komplexe Zahl ist.

Auf jeden Fall können wir aber nun jede Gleichung der Form $w^2 = z$ lösen und wir können sogar noch mehr.

Sind a, b und c beliebige komplexe Zahlen mit $a \neq 0$, dann hat die quadratische Gleichung $aw^2 + bw + c = 0$ immer mindestens eine Lösung (und höchstens zwei verschiedene).

Satz 14.4

Im Prinzip geht man dabei wie beim Lösen solcher Gleichungen in der Schule vor, z. B. mithilfe der quadratischen Ergänzung oder der sogenannten „ p - q -Formel“.²⁸ Wenn zwischendurch eine Wurzel auftaucht, wendet man Satz 14.3 bzw. die im Anschluss vorgestellte Methode an.

Um etwa $2w^2 + (2 + 2i)w + (12 + 36i) = 0$ zu lösen, dividiert man zunächst durch 2 und erhält $w^2 + (1 + i)w + (6 + 18i) = 0$. Will man die p - q -Formel anwenden, so hat man $p = 1 + i$ und $q = 6 + 18i$. $p^2/4 - q$ ist $-6 - 35i/2$. Eine „Wurzel“ dieser Zahl ist $5/2 - 7i/2$. Kombiniert mit $-p/2 = -1/2 - i/2$ ergeben sich die zwei Lösungen $2 - 4i$ und $-3 + 3i$. (Machen Sie eine Probe!)

Aufgabe 14.20. Wie in Aufgabe 14.19 sollten Sie sich selbst Aufgaben stellen und diese dann später lösen. Aufgaben generiert man, indem man sich zwei komplexe Zahlen w_1 und w_2 (die Lösungen) ausdenkt und dann den Term $(w - w_1)(w - w_2)$ ausmultipliziert.

Aufgabe 14.21. Wie viele Lösungen hat die Gleichung $x^4 = 16$ und wie sehen diese aus? (Hinweis: Machen Sie eine Substitution wie oben.)

★Aufgabe 14.22. Für den Fall, dass Sie den Dingen gerne auf den Grund gehen: In Lemma 14.1 ist eine Aussage versteckt, die nicht völlig selbstverständlich ist. Dass $z = a_1 + b_1i$ und $w = a_2 + b_2i$ nur dann identisch sein können, wenn $a_1 = a_2$ und $b_1 = b_2$ gilt, folgt daraus, dass die Differenz $(a_1 - a_2) + (b_1 - b_2)i$ „offenbar“ nur dann verschwindet, wenn $a_1 - a_2 = 0$ und $b_1 - b_2 = 0$ gilt. Aber wieso ist das klar? Könnte es nicht sein, dass es reelle Zahlen a und b mit $a + bi = 0$ und $\{a, b\} \neq \{0\}$ gibt?

Die komplexen Zahlen bilden gewissermaßen den Abschluss der Erweiterung der Zahlenwelt, weil man mit Ihnen *alle* Gleichungen lösen kann, die sich mit einer

²⁸Siehe dazu mein [Video zur geometrischen Herleitung der quadratischen Ergänzung](#).

Variablen, komplexen Zahlen sowie Addition, Subtraktion und Multiplikation hinschreiben lassen – also z. B. $x^4 + 4x^3 + x^2 + 24x + 260 = 0$. (Siehe auch Aufgabe 14.21.) Mehr dazu im [Video über den sogenannten Fundamentalsatz der Algebra](#).

Bemerkungen

Für komplexe Zahlen werden üblicherweise die Variablennamen z und w (statt x und y) verwendet. Aber natürlich ist das keine feste Regel, sondern nur eine Konvention.

Beachten Sie, dass komplexe Zahlen aus zwei Teilen bestehen, die zusammengehören, und dass deshalb oft Klammern nötig sind. Für $z = 3 + 4i$ sind $5 \cdot z$ und $5 \cdot 3 + 4i$ beispielsweise *nicht* dasselbe. $3/4 + 5i/4$ und $3 + 5i/4$ sind ebenfalls zwei unterschiedliche Zahlen. (Siehe auch die Bemerkungen zu Klammern am Ende von Abschnitt 5.)

Obwohl komplexe Zahlen immer zwei Teile haben, werden die Teile, die verschwinden, natürlich in der Regel nicht hingeschrieben. Statt $3 + 0i$ schreibt man einfach 3 und statt $0 + 4i$ einfach $4i$.

Vermeiden Sie die typischen Fehler beim Umgang mit dem Imaginärteil, auf die in Aufgabe 14.2 hingewiesen wurde.

★ Komplexe Zahlen im Computer

PYTHON bietet als Standard bereits komplexe Zahlen. Fügt man direkt hinter einer Zeichenkette, die als Zahl erkennbar ist, den Buchstaben j (klein oder groß) an, so wird das als rein imaginäre Zahl erkannt.²⁹ Ein alleinstehendes j wird jedoch als die Variable mit diesem Namen interpretiert. Beliebige komplexe Zahlen gibt man als Summe von Real- und Imaginärteil ein oder mit der Funktion `complex`.

```
42J, 42j, 2.3j, 23e-1j
2+3j, complex(2,3)
(5+3j)*(3+2j)
z = 5+3j
abs(z), z.real, z.imag
```

Das letzte Beispiel zeigt, dass Real- und Imaginärteil intern grundsätzlich als Fließkommazahlen gespeichert werden.³⁰ Für das Rechnen mit komplexen Zahlen sollte in der Regel die Standardbibliothek `CMATH` benutzt werden.

```
import cmath
# würde mit dem "normalen" sqrt nicht funktionieren:
```

²⁹PYTHON übernimmt hier die in den Ingenieurwissenschaften übliche Bezeichnung j für die imaginäre Einheit.

³⁰Sprachen wie [JULIA](#) oder [COMMON LISP](#) bieten auch [Datentypen](#) an, in denen die beiden Teile einer komplexen Zahl intern als ganze Zahl oder als Bruch gespeichert werden können.


```
cmath.sqrt(-15-8j), cmath.sqrt(-4)  
cmath.phase(-42)           # das Argument
```


IV

Funktionen

Der Begriff *Funktion* kommt bereits in der Schule regelmäßig vor, weil es sich um ein zentrales Konzept der Mathematik handelt. Allerdings ist die Vorstellung von Funktionen, die man aus der Schule mitbringt, häufig zu eng. In diesem Kapitel erarbeiten wir uns die moderne mathematische Sichtweise.

15. Tupel

Wie haben bereits **Mengen** als Grundbausteine der Mathematik kennengelernt und regelmäßig genutzt. Mengen haben zwei Eigenschaften, die oft hilfreich sind, manchmal aber auch unpraktisch: Sie kennen keine Reihenfolge und es gibt keine Dopplungen. $\{1, 4, 7\}$, $\{7, 1, 4\}$ und $\{1, 7, 4, 4\}$ sind drei verschiedene Arten, *dieselbe* Menge aufzuschreiben. Wir führen deshalb eine weitere Sorte von „Containern“ ein, die diese Eigenschaften nicht haben:

Sind $a_1, a_2, \dots, a_n$ irgendwelche (mathematischen) Objekte (die nicht notwendig verschieden sein müssen), dann bezeichnet man deren Zusammenfassung $(a_1, a_2, \dots, a_n)$ als **n -Tupel** mit den **Komponenten** a_1 bis a_n . Die definierende Eigenschaft von Tupeln ist, dass zwei Tupel $(a_1, \dots, a_m)$ und $(b_1, \dots, b_n)$ dann und nur dann gleich sind, wenn $m = n$ und $a_1 = b_1$, $a_2 = b_2$ und so weiter bis $a_m = b_m$ gilt.

2-Tupel nennt man auch **(geordnete) Paare**. Zu 3-Tupeln sagt man manchmal auch **Tripel** und es gibt analoge Bezeichnungen für Tupel mit mehr als drei Komponenten wie **Quadrupel**, **Quintupel** et cetera.

Der Vollständigkeit halber setzt man noch fest, dass $()$ das *einzige* 0-Tupel ist, das man auch **das leere Tupel** nennt.

Definition

Nach dieser Definition sind $(1, 2)$ und $(2, 1)$ zwei verschiedene geordnete Paare, wohingegen $\{1, 2\}$ dieselbe Menge wie $\{2, 1\}$ ist. Und (2) und $(2, 2)$ sind zwei verschiedene Tupel (ein 1-Tupel und ein 2-Tupel), weil sie eine unterschiedliche Anzahl von Komponenten haben. $\{2\}$ und $\{2, 2\}$ unterscheiden sich hingegen nicht. Insbe-

sondere sind 42, {42} und (42) drei *verschiedene* Objekte: eine Zahl, ein **Singleton** (eine Menge mit einem Element) und ein 1-Tupel.

Weil Tupel eine Reihenfolge haben, kann man auch von der *ersten* oder *siebten* Komponente eines Tupels sprechen, während ein Ausdruck wie „das dritte Element“ bei Mengen keinen Sinn ergibt. Aus der Sicht der Informatik kann man sich Tupel wie **Arrays** oder **Listen** vorstellen.

Wir werden zunächst hauptsächlich mit Paaren arbeiten, also mit 2-Tupeln.

Definition

Das **kartesische Produkt** oder **Mengenprodukt** der Mengen A und B ist die Menge $A \times B = \{(a, b) : a \in A, b \in B\}$. Für $A \times A$ schreibt man auch A^2 .

$A \times B$ ist also die Menge aller Paare, bei denen die erste Komponente ein Element von A und die zweite eines von B ist. Mengenprodukte kann man sich ganz gut in Tabellenform vorstellen. Mit den Mengen $A = \{1, 10, 17, 23\}$ und $B = \{10, 23, 42\}$ sieht $A \times B$ so aus:

	10	23	42
1	(1, 10)	(1, 23)	(1, 42)
10	(10, 10)	(10, 23)	(10, 42)
17	(17, 10)	(17, 23)	(17, 42)
23	(23, 10)	(23, 23)	(23, 42)

Man beachte, dass nach Definition des geordneten Paares alle Einträge in dieser Tabelle zwangsläufig verschieden sein müssen. Es kann keine Dopplungen geben. Deshalb können wir sofort etwas über die Mächtigkeit von $A \times B$ aussagen, was auch die Schreibweise als Produkt rechtfertigt.

Lemma 15.1

Sind A und B endliche Mengen, dann gilt $|A \times B| = |A| \cdot |B|$.

Aber wie sieht es mit unendlichen Mengen aus? Wie kann man sich deren kartesisches Produkt vorstellen? Dazu zwei wichtige Beispiele.

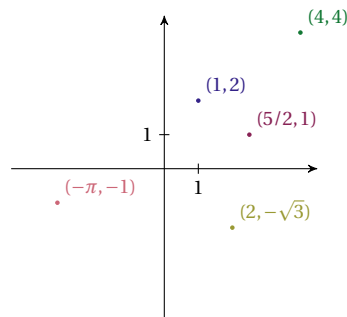
Ein Produkt wie $\mathbb{N} \times \mathbb{N}$ ist auch nichts anderes als eine Tabelle, aber eine mit unendlich vielen Zeilen und Spalten, die jeweils mit den Zahlen 0, 1, 2 und so weiter durchnummeriert sind. Es geht nach rechts und nach unten jeweils – durch Auslassungspunkte angedeutet – immer weiter.

	0	1	2	3	...
0	(0, 0)	(0, 1)	(0, 2)	(0, 3)	...
1	(1, 0)	(1, 1)	(1, 2)	(1, 3)	...
2	(2, 0)	(2, 1)	(2, 2)	(2, 3)	...
3	(3, 0)	(3, 1)	(3, 2)	(3, 3)	...
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$

Das Produkt $\mathbb{R}^2$ kennen Sie schon aus der Schule, obwohl Sie es bisher vielleicht nicht unter diesem Blickwinkel betrachtet haben. Man stellt sich dafür die Menge $\mathbb{R}$ der reellen Zahlen als **Zahlengerade** vor, legt zwei solche Geraden senkrecht

zueinander so auf die Ebene, dass sie sich an der Stelle treffen, wo jeweils die Null sitzt, und erhält dadurch eine eindeutige Zuordnung zwischen den Punkten der Ebene und deren „Adressen“ (Koordinaten), die geordnete Paare reeller Zahlen sind. Das ist ein sogenanntes *kartesisches Koordinatensystem*.

kartesisches
Koordinatensystem



Stellen Sie sich nun vor, wo die Elemente der Menge $\mathbb{N} \times \mathbb{N}$ in diesem Koordinatensystem zu finden sind. Sie erhalten ein *Gitter* aus Punkten.

Aufgabe 15.1. Die Menge B aus Aufgabe 6.5 könnte man auch als

$$B = \{mn : (m, n) \in A^2\}$$

schreiben. Machen Sie sich klar, warum B weniger Elemente als A^2 hat.

Aufgabe 15.2. Sei $M = \{23, 42, \pi\}$. Wie sehen $\{0\} \times M$ und $\emptyset \times M$ aus?

Aufgabe 15.3. Ist das Mengenprodukt kommutativ, gilt also $A \times B = B \times A$ für alle Mengen A und B ? Begründen Sie Ihre Antwort.

Aufgabe 15.4. Können Sie eine Menge A angeben, für die $|A \times A| = 42$ gilt? Können Sie Mengen A und B angeben, für die $|A \times B| = 42$ gilt?

Aufgabe 15.5. Sei $A = [6]$, $B = \mathcal{P}(A)$, $C = B \times B$ und $D = C \setminus \{\emptyset\}$. Geben Sie $|D|$ an.

Aufgabe 15.6. Sei $A_q = \{p \in \mathbb{P} : p < q\}$. Wie sehen die Mengen A_4 und A_5 aus? Für welche Primzahl q gilt $|A_q \times A_q| = 100$?

★Aufgabe 15.7. Sei $A = \{1\}$. Gilt dann $A \times (A \times A) = (A \times A) \times A$?

Wir verallgemeinern nun das Mengenprodukt A^2 .

Ist A eine Menge und n eine natürliche Zahl, so bezeichnen wir mit A^n die Menge aller n -Tupel mit Komponenten aus A :

$$A^n = \{(a_1, \dots, a_n) : \forall i \in [n] \ a_i \in A\}$$

Definition

Insbesondere ist damit A^0 *immer* die einelementige Menge $\{()\}$. Die folgende Aussage sollte nach dieser Definition offensichtlich sein.

Lemma 15.2

Ist A endlich, dann gilt $|A^n| = |A|^n$.

Aufgabe 15.8. Höchstwahrscheinlich haben Sie sich irgendwann in der Schule auch mit **Geometrie** im dreidimensionalen Raum beschäftigt. Überlegen Sie sich (analog zum Beispiel $\mathbb{R}^2$ oben), wie man sich die Mengen $\mathbb{R}^3$ vorstellen kann.

Bemerkungen

Obwohl in den einführenden Beispielen meistens Zahlen vorkommen, sollten Sie nicht vergessen, dass die Komponenten von Tupeln (genau wie die Elemente von Mengen) *irgendwelche* mathematischen Objekte sein können. Dazu gehören neben Zahlen eben auch Mengen, Tupel oder Dinge, die in diesem Skript bisher noch gar nicht vorkamen, wie beispielsweise Matrizen, Vektoren, **Funktionen** und so weiter.

Ihnen ist vielleicht aufgefallen, dass mit einem Ausdruck wie $(3, 4)$ sowohl ein Tupel als auch ein **Intervall** gemeint sein kann. Das ist zwar etwas unglücklich, aber es lässt sich nicht vermeiden, dass man in der Mathematik ab und zu dieselbe Bezeichnung für verschiedene Dinge verwendet. Es sollte in der Regel aus dem Kontext hervorgehen, was gemeint ist.

Ein häufiges Missverständniss ist, dass die Menge A^2 etwas mit dem Quadrieren von Zahlen zu tun hat. Dem ist jedoch nicht so. A^2 ist lediglich – wie in der Definition des kartesischen Produktes vereinbart – eine Abkürzung für $A \times A$.

Weil das gerne vergessen wird, sei noch einmal darauf hingewiesen, dass 42 , $\{42\}$ und (42) drei verschiedene Objekte sind.

★ Tupel im Computer

Auch in PYTHON gibt man **Tupel** mit Klammern und Kommata ein.¹ Lediglich bei 1-Tupeln muss man ein „überflüssiges“ Komma hinzufügen.

```
(42,23,1) == (1,23,42)    # vergleiche mit {42,23,1} == {1,23,42}
type((42,))               # vergleiche mit type((42))
```

Das Mengenprodukt zweier Mengen könnte man so definieren:

```
product = lambda A,B: {(a,b) for a in A for b in B}
```

Die Standardbibliothek **ITERTOOLS** bietet allerdings eine flexiblere und effizientere Funktion `product`.

¹In vielerlei Hinsicht sind die Tupel von PYTHON das „gefrorene“ Pendant der **Listen**.

16. Funktionen

Das Wort *Funktion* stammt von [Gottfried Wilhelm Leibniz](#) aus dem 17. Jahrhundert. Daran erkennt man, dass dieses Konzept wesentlich jünger als andere Bereiche der Mathematik ist. Bis zum heutigen Funktionsbegriff war es allerdings ein weiter Weg. Noch bis ins 19. Jahrhundert hinein bestanden nur recht schwammige Vorstellungen davon, was genau als Funktion gelten sollte, und sie waren eng mit dem Konzept einer *Rechenvorschrift* verbunden. Erst durch diverse [pathologische Beispiele](#), die u.a. von [Bernard Bolzano](#), [Peter Gustav Lejeune Dirichlet](#) und [Karl Weierstraß](#) entwickelt wurden, wurde nach und nach klar, dass man ein allgemeineres Konzept benötigte. Die inzwischen übliche Definition entstand erst im 20. Jahrhundert und basiert auf der [Mengenlehre](#).

Auch in der Schule verbindet man mit dem Funktionsbegriff im Allgemeinen eine Rechenvorschrift, die Zahlen andere Zahlen zuordnet. Ein „Klassiker“ ist die Wurzelfunktion, die einer Zahl x ihre Quadratwurzel $\sqrt{x}$ zuordnet. Ein anderes Beispiel ist eine Gerade mit der Steigung m und dem Achsenabschnitt b , die einer Zahl x die Zahl $mx + b$ zuordnet.²

Anhand solcher Beispiele wollen wir herausarbeiten, was alle Funktionen im mathematischen Sinne gemeinsam haben, inwieweit diese Beispiele jedoch zu speziell sind und nicht den gesamten Funktionsbegriff umfassen:

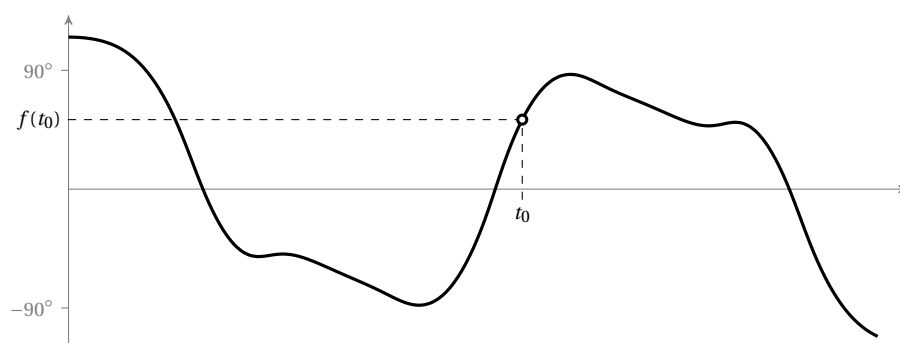
- (i) Man kann sich Funktionen als „Vorrichtungen“ oder „Maschinen“ vorstellen, die eine Eingabe erhalten und dazu eine Ausgabe liefern. Die Wurzelfunktion liefert z. B. bei der Eingabe 4 die Ausgabe 2 und bei der Eingabe 9 die Ausgabe 3. In diesem Sinne haben Funktionen große Ähnlichkeit mit [Funktionen in Programmiersprachen](#).
- (ii) Mathematische Funktionen sind „verlässlich“: Für dieselbe Eingabe liefern sie immer dieselbe Ausgabe.³ Das unterscheidet sie von Funktionen in Programmiersprachen, die u.a. auch [zufällige Werte](#) ausgeben, je nach Uhrzeit unterschiedliche Dinge tun oder ein „Gedächtnis“ haben können.
- (iii) Nicht jeder Eingabewert ist zulässig. Die erwähnte Wurzelfunktion akzeptiert etwa keine negativen Zahlen als Eingabe. Das kennen Sie aus der Schule als [Definitionsbereich](#) der Funktion.
- (iv) Es muss nicht um Zahlen gehen. Sowohl die Eingabewerte als auch die Ausgabewerte von Funktionen können beliebige mathematische Objekte wie Mengen, Tupel, Vektoren, Matrizen und so weiter sein. Wir werden dafür schon bald Beispiele sehen.
- (v) Anders als in den genannten Beispielen muss es keine *Rechenvorschrift* geben, die festlegt, *wie* die Ausgabe aus der Eingabe generiert wird. Es geht einzig und allein darum, *dass* es zu jeder möglichen Eingabe eine klar defi-

²In der Schule wird so etwas oft als „lineare Funktion“ bezeichnet, was sich aber leider mit der Hochschulmathematik beißt. Besser wäre das Adjektiv *affin-linear*.

³Man kann es auch so sehen, dass Mathematik „zeitlos“ ist, was teilweise schon im [ersten Kapitel](#) thematisiert wurde. Es ist egal, *wann* man die Funktion mit einer Eingabe füttert.

nierte Ausgabe gibt. (Beispielsweise könnte diese einer vorher festgelegten Tabelle entstammen, die „ausgewürfelt“ wurde.)

Besonders der letzte Punkt ist wohl erklärungsbedürftig. Warum belässt man es nicht einfach bei Funktionen, die man berechnen kann? Eine Antwort darauf liefert die Physik. Dort ist man häufig auf der Suche nach Funktionen – meistens nach solchen, die bestimmte Werte in Abhängigkeit von der Zeit liefern: Welche Position hat ein Himmelskörper zu einem bestimmten Zeitpunkt? Welche Geschwindigkeit hat ein Elementarteilchen zu einem bestimmten Zeitpunkt? Und in der Regel kann die Physik mit sogenannten **Differentialgleichungen** die gesuchten Funktionen so beschreiben, dass man mathematisch beweisen kann, dass sie existieren und eindeutig bestimmt sind. Gleichzeitig kann man aber auch beweisen, dass man die Funktionswerte nicht exakt berechnen kann. Man kann mit anderen Worten kein Computerprogramm schreiben, das sie auswirft, indem es sie durch eine Formel berechnet. Die folgende Grafik zeigt ein Beispiel für so eine Funktion: den zeitlichen Verlauf eines der Winkel eines **Doppelpendels**.



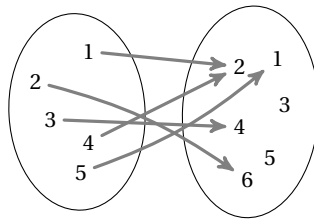
Die grafische Darstellung wurde zwar durch einen Computer erzeugt, aber die dafür benötigten Zahlen konnten nur näherungsweise durch eine Simulation ermittelt werden. Wir sehen die Visualisierung der Funktion und wir können uns vorstellen, wie wir für einen Zeitpunkt t_0 den Winkel $f(t_0)$ ablesen. Aber es ist eben nur so wie das Ablesen der Zahlen aus einer Tabelle.⁴

Die moderne Vorstellung einer Funktion ist daher auch die einer Tabelle mit zwei Spalten: einer für die Eingabe und einer für die Ausgabe.

Eingabe	Ausgabe
1	2
2	6
3	4
4	2
5	1

⁴Es hätte auch so sein können, dass die Punkte, aus denen die Grafik zusammengesetzt ist, Messwerte sind. Die hätten dann ja ebenfalls einmal in einer Tabelle gestanden.

Manchmal ist es auch hilfreich, sich eine Funktion mittels eines Pfeildiagramms zu visualisieren:⁵



Da die Reihenfolge der Zeilen nicht relevant ist, kann man die Tabelle auch durch eine Menge von Paaren repräsentieren:

$$R = \{(1, 2), (2, 6), (3, 4), (4, 2), (5, 1)\} \quad (16.1)$$

R können wir beispielsweise entnehmen, dass der Eingabe 3 die Ausgabe 4 zugeordnet wird. Und wir können an R auch sehen, dass als Eingaben nur die natürlichen Zahlen von 1 bis 5 zugelassen sind. Aus mathematischer Sicht *ist* R eine Funktion, weil man mehr über eine Funktion nicht wissen muss.

Aufgabe 16.1. Sie haben gerade gelesen, dass die Reihenfolge der Zeilen nicht relevant ist und man daher die Tabelle als Menge (die ja keine Reihenfolge kennt) darstellen kann. Wieso ist die Reihenfolge denn unwichtig?

Bevor wir durch eine Definition festhalten, was genau eine Funktion ist, müssen wir jedoch noch klären, dass nicht *jede* Menge von Paaren sich als Funktion eignet. Wenn wir eine kleine Änderung vornehmen und statt R die Menge

$$S = \{(1, 2), (2, 6), (3, 4), (4, 2), (4, 1)\}$$

betrachten, dann ist nicht mehr klar, was bei der Eingabe 4 geschehen soll: Soll 2 oder 1 ausgegeben werden? So etwas darf nicht passieren. Das halten wir jetzt fest:

Eine Menge von Paaren nennt man eine **Relation**. Eine Relation R wird **rechtseindeutig** genannt, wenn aus $(x, y_1) \in R$ und $(x, y_2) \in R$ immer $y_1 = y_2$ folgt.⁶ Eine **Funktion** ist eine rechtseindeutige Relation.

Definition

Funktionen werden oft auch *Abbildungen* genannt.⁷

Aufgabe 16.2. Was bedeutet Rechtseindeutigkeit für Pfeildiagramme wie das obige? Was darf man in so einem Diagramm also nicht sehen?

⁵Das ist allerdings *kein* Mengendiagramm im Sinne des [Abschnitts über Mengen](#), denn hier werden Objekte ja mehrfach gezeichnet. Bis auf die Sechs tauchen alle Zahlen links und rechts auf.

⁶ R darf also keine zwei Paare (x_1, y_1) und (x_2, y_2) mit $x_1 = x_2$ und $y_1 \neq y_2$ enthalten. Im Englischen wird so eine Relation auch *univalent* genannt.

⁷In manchen Büchern wird aus didaktischen Gründen zwischen Funktionen und Abbildungen unterschieden. Wir werden das aber nicht tun. Es ist allerdings so, dass bestimmte Funktionen aus historischen Gründen eher als Abbildungen bezeichnet werden, z. B. die linearen Abbildungen der linearen Algebra.

Aufgabe 16.3. Welche der folgenden Mengen ist eine Funktion?

$$f_1 = \{(n, n/3) : n \in \mathbb{N}\}$$

$$f_2 = \{(x, y) \in \mathbb{R}^2 : x^2 + y^2 = 1\}$$

$$f_3 = \{(1, 2), 3, (4, 5)\}$$

$$f_4 = \{(1, 3), (3, 5), (4, 5)\}$$

$$f_5 = \{(1, 2), (1, 3), (4, 5)\}$$

$$f_6 = \emptyset$$

Definition

Ist f eine Relation, dann nennen wir die Menge $\text{dom}(f) = \{x : (x, y) \in f\}$ den **Definitionsbereich** der Relation und $\text{rng}(f) = \{y : (x, y) \in f\}$ ihren **Wertebereich** bzw. ihre **Bildmenge**.⁸

Der Definitionsbereich ist also einfach die Menge aller ersten Komponenten der Paare der Relation und der Wertebereich die Menge aller zweiten Komponenten. In einem Pfeildiagramm wie auf Seite 125 ist der Definitionsbereich die Menge der Objekte, von denen ein Pfeil ausgeht, und der Wertebereich die Menge der Objekte, auf die mindestens ein Pfeil zeigt. Beachten Sie, dass nach unserer Definition jede Relation (und damit insbesondere jede Funktion) ihren Definitions- und Wertebereich „mit sich herumträgt“: die beiden Mengen sind fest in die Relation „hineincodiert“.

Aufgabe 16.4. Geben Sie Definitions- und Wertebereich der Relationen an, die in Aufgabe 16.3 vorkamen.

Aufgabe 16.5. Wenn f eine endliche Funktion ist, also eine Menge, die nur endlich viele Paare enthält, welche Zusammenhänge bestehen dann zwischen $|f|$, $|\text{dom}(f)|$ und $|\text{rng}(f)|$?

Die Menge f_1 aus Aufgabe 16.3 ist das „klassische“ Beispiel für eine Funktion. Hier wird jeder Eingabe n eine Ausgabe durch die Rechenvorschrift $n/3$ zugeordnet. Solche Funktionen schreibt man ganz ausführlich in der folgenden Form:

$$f_1 : \begin{cases} \mathbb{N} \rightarrow \mathbb{Q} \\ n \mapsto n/3 \end{cases} \quad (16.2)$$

Zielmenge

Vor dem Doppelpunkt steht optional der Name der Funktion. In der oberen Zeile steht links der Definitionsbereich und rechts eine sogenannte *Zielmenge*. In der unteren Zeile steht links ein Platzhalter für ein Element des Definitionsbereichs und rechts die Rechenvorschrift angewandt auf diesen Platzhalter. (Linke und rechte Seite der unteren Zeile sind also die beiden Komponenten eines Paares der Funktion, wenn man sie als Relation betrachtet.) Beachten Sie, dass die beiden Pfeile unterschiedlich aussehen.⁹

Verwechseln Sie die Zielmenge (engl. *codomain*) nicht mit dem Wertebereich!¹⁰ Der Wertebereich ist eine durch die Funktion eindeutig bestimmte Menge, während die

D_f W_f

⁸Wir verwenden dabei die internationalen Bezeichnungen *domain* und *range*. In der Schule schreibt man für den Definitionsbereich von f oft auch D_f und für den Wertebereich W_f .

⁹Der untere Pfeil wird *maplet* genannt.

¹⁰Sie wussten aus Aufgabe 16.4 bereits, dass $\mathbb{Q}$ nicht der Wertebereich von f_1 ist.

Zielmenge irgendeine Obermenge des Wertebereichs ist. Wir hätten in (16.2) statt $\mathbb{Q}$ beispielsweise auch $\mathbb{R}$ schreiben können und hätten damit *dieselbe* Funktion definiert!¹¹

Aufgabe 16.6. Bei der Angabe der Zielmenge kann man also irgendeine Obermenge des Wertebereichs angeben. Warum ist man da so „liberal“? Es wäre doch präziser, wenn man in der Schreibweise (16.2) statt der Zielmenge den Wertebereich angeben würde. Haben Sie eine Idee, warum man es so macht?

Ist f eine Funktion mit Definitionsbereich A und Zielmenge B , dann sagt man auch, dass f eine Funktion *von A nach B* ist und schreibt als Abkürzung dafür $f: A \rightarrow B$. Ebenso sagt man dann, dass f *auf A* definiert ist.

$$f: A \rightarrow B$$

Durch die Rechtseindeutigkeit ist sichergestellt, dass es für jedes Element x des Definitionsbereichs einer Funktion f genau ein y mit $(x, y) \in f$ gibt. Für dieses eindeutig bestimmte y schreibt man $f(x)$ und nennt es das *Bild von x unter f* . Man sagt auch, dass die Funktion x *auf $f(x)$ abbildet* und schreibt dafür auch $x \mapsto f(x)$. x und $f(x)$ entsprechen also dem, was wir am Anfang des Abschnitts Ein- und Ausgabe genannt haben. Weitere verbreitete Bezeichnungen sind *Argument* für x und *Funktionswert* für $f(x)$.

$$f(x)$$

$$x \mapsto f(x)$$

Argument

Funktionswert

Manchmal werden auch informelle Sprechweisen wie „die Funktion $x \mapsto x^2$ “ oder „die Funktion $f(x) = x^2$ “ verwendet. Das sollte jedoch nur geschehen, wenn aus dem Kontext der Definitionsbereich der Funktion hervorgeht, weil sie nur dann vollständig beschrieben ist. In diesem Sinne kann man die im Folgenden definierte Funktion auch durch $x \mapsto x$ beschreiben.¹²

Für jede Menge A definieren wir die **Identität** auf A als die Menge

$$\text{id}_A = \Delta_A = \{(x, y) \in A^2 : x = y\} = \{(x, x) : x \in A\}.$$

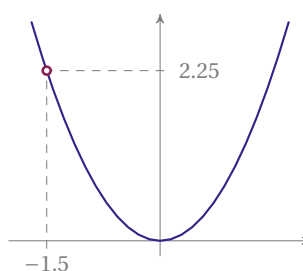
Definition

Da wir Funktionen als Mengen von Paaren betrachten, kann man Funktionen, die reelle Zahlen auf reelle Zahlen abbilden (bei denen also sowohl Definitions- als auch Wertebereich Teilmengen von $\mathbb{R}$ sind), in einem **Koordinatensystem** darstellen. Das kennen Sie aus der Schule als *Funktionsgraph*. Die folgenden Grafik zeigt z. B. einen Ausschnitt des Funktionsgraphen der Funktion $x \mapsto x^2$.

Funktionsgraph

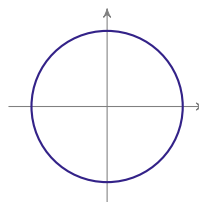
¹¹Auch hier muss ich leider wieder darauf hinweisen, dass in der Mathematik nicht immer Einigkeit besteht. Es gibt auch Bücher, in denen die Zielmenge fester Bestandteil der Funktionsdefinition ist. In diesem Fall würden wir beim Wechsel von $\mathbb{Q}$ nach $\mathbb{R}$ tatsächlich eine *andere* Funktion erhalten. In solchen Büchern ist eine Funktion aber auch nicht mehr einfach eine Menge von Paaren.

¹²Wobei allerdings zu beachten ist, dass es nicht um eine Funktion geht, sondern man für jede Menge A eine andere Funktion erhält.



Beachten Sie, dass der exemplarisch hervorgehobene Punkt mit den Koordinaten $(-1.5, 2.25)$ nicht nur dafür steht, dass 2.25 der Funktionswert von -1.5 ist, sondern dass das Paar $(-1.5, 2.25)$ nach unserer Definition tatsächlich ein Element der Funktion ist. Der Funktionsgraph *ist* also die Funktion (bzw. ihre grafische Repräsentation).

Natürlich kann man auch beliebige Relationen, die Teilmengen von $\mathbb{R} \times \mathbb{R}$ sind, auf dieselbe Art visualisieren. Die folgende Grafik zeigt beispielsweise den Einheitskreis (siehe Aufgabe 16.3).



Aufgabe 16.7. Wie kann man an so einer Grafik erkennen, dass die Relation nicht rechtseindeutig, also keine Funktion, ist?

Eine der angenehmsten Eigenschaften, die Funktionen haben können, ist die Injektivität. Dann kann man sie nämlich *umkehren*.

Definition

Eine Funktion f wird **injektiv** genannt, wenn für alle $x, y \in \text{dom}(f)$ aus $x \neq y$ immer $f(x) \neq f(y)$ folgt.

Injektivität bedeutet in Worten, dass für unterschiedliche Eingaben (Argumente) auch unterschiedliche Ausgaben (Funktionswerte) herauskommen. Am Pfeildia-gramm von Seite 125 kann man das ganz gut nachvollziehen. Die dort dargestellte Funktion ist nicht injektiv, denn die Zahl 2 wird von zwei Pfeilen getroffen.¹³

Ein „klassisches“ Beispiel, an dem man sich dieses neue Konzept ebenfalls gut vergegenwärtigen kann, ist die Funktion

$$f: \begin{cases} \mathbb{R} \rightarrow \mathbb{R} \\ x \mapsto x^2 \end{cases} . \quad (16.3)$$

linkseindeutig

¹³Erinnern Sie sich, dass Rechtseindeutigkeit bedeutet, dass von jedem Objekt links nur ein Pfeil ausgeht. Injektivität bedeutet, dass bei jedem Objekt rechts (höchstens) ein Pfeil ankommt. Man nennt Injektivität daher manchmal auch *Linkseindeutigkeit*.

Sie ist *nicht* injektiv, denn beispielsweise gilt $f(2) = 4 = f(-2)$, obwohl 2 und -2 natürlich zwei verschiedene Argumente sind.

Aufgabe 16.8. Fällt Ihnen eine ganz einfache, aus der Schule bekannte Klasse von Funktionen von $\mathbb{R}$ nach $\mathbb{R}$ ein, die alle injektiv sind?

Aufgabe 16.9. Im Lösungshinweis für Aufgabe 16.8 habe ich absichtlich etwas unpräzise formuliert. Ist Ihnen aufgefallen, was man noch hätte hinzufügen müssen?

Aufgabe 16.10. Machen Sie sich klar, dass eine Funktion f genau dann injektiv ist, wenn aus $f(x) = f(y)$ immer $x = y$ folgt.

Aufgabe 16.11. Wie kann man am Funktionsgraphen erkennen, ob eine Funktion injektiv ist? (Tipp: Siehe Aufgabe 16.7.)

Aufgabe 16.12. Die folgenden Funktionen sind alle nicht injektiv. Begründen Sie das jeweils mit konkreten Gegenbeispielen.

$$f = \{(n, n \bmod 42) : n \in \mathbb{N}\}$$

$$g = \{(A, |A|) : A \subseteq [6]\}$$

$$h = \{(n, Q(n)) : n \in \mathbb{N}\}$$

Dabei soll Q die *iterierte Quersumme* sein.

Die Bedeutung der Injektivität steckt in der folgenden Aussage, die ich nicht formal beweisen werde, die man sich aber leicht anhand eines Pfeildiagramms klarmachen kann.

Eine Funktion f ist genau dann injektiv, wenn die Relation

$$f^{-1} = \{(y, x) : (x, y) \in f\}$$

eine Funktion ist (die man dann die *Umkehrfunktion*) von f nennt).

Satz 16.1

Injektive Funktionen sind also die, die man *umkehren* kann. (In f^{-1} tauschen die Komponenten jedes Paares von f einfach ihre Rollen.) Jede Funktion ordnet jedem Argument eindeutig einen Funktionswert zu. Injektive Funktionen ordnen zusätzlich auch jedem Funktionswert eindeutig ein Argument zu. Häufig kann man – wie in Aufgabe 16.14 – auch eine Rechengvorschrift für die Umkehrfunktion angeben, die dann gleichzeitig ein Beleg dafür ist, dass man es tatsächlich mit einer injektiven Funktion zu tun hat. Mit anderen Worten: Wenn Sie ein Verfahren angeben können, wie man aus jeder Ausgabe *eindeutig* die zugehörige Eingabe rekonstruieren kann, dann ist die Funktion injektiv.

Aufgabe 16.13. Wenn f injektiv ist, was ist dann der Definitionsbereich der Umkehrfunktion f^{-1} ?

Aufgabe 16.14. Begründen Sie für die folgenden Funktionen, dass sie injektiv sind, indem Sie jeweils eine Rechenvorschrift für die Umkehrfunktion angeben.

$$f = \{(x, 2x + 8) : x \in \mathbb{R}\}$$

$$g = \{(A, [6] \setminus A) : A \subseteq [6]\}$$

$$h = \{(n, n^2) : n \in \mathbb{N}\}$$

Aufgabe 16.15. In Aufgabe 16.14 kam eine Funktion vor, die ihre eigene Umkehrfunktion war. Können Sie noch ein (ganz einfaches) Beispiel für so eine Funktion angeben?

Aufgabe 16.16. Sei $A = \{x \in \mathbb{R} : x \geq 0\}$ und f die Funktion, die jedes $x \in \mathbb{R}$ auf den Absolutbetrag $|x|$ abbildet. Ist f injektiv?

Im Folgenden definieren wir keinen neuen Begriff, aber wir lernen eine weitere Art kennen, wie in der Praxis Funktionen eingesetzt werden. Nämlich die, dass die Argumente der Funktion **Tupel** sind. Das ist nichts Neues und durch die Definitionen der vorherigen Seiten bereits abgedeckt, aber es wird häufig anders notiert. Hier ein Beispiel:

$$f: \begin{cases} \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N} \\ (m, n) \mapsto m^2 + n^2 \end{cases} \quad (16.4)$$

zweistellige Funktion
Verknüpfung

Diese Funktion bildet z. B. das Argument $(1, 5)$ auf $1^2 + 5^2 = 26$ ab und das Argument $(4, 4)$ auf $4^2 + 4^2 = 32$. Man müsste also eigentlich $f((1, 5)) = 26$ schreiben. In der Regel wird man sich jedoch die „überflüssigen“ Klammern sparen und einfach $f(1, 5)$ schreiben. Man spricht dann auch von einer *zweistelligen* Funktion¹⁴ und verwendet manchmal auch das Wort *Verknüpfung* statt Funktion. Wenn Sie schon ein bisschen programmiert haben, wird Ihnen das nicht neu vorkommen. Man beachte jedoch die Feinheiten:

- Der Prototyp des Arguments links vom Zeichen $\mapsto$ muss zu den Elementen des Definitionsbereichs passen. Im Beispiel ist (m, n) in Ordnung, weil die Elemente von $\mathbb{N} \times \mathbb{N}$ Paare sind. (m, m) wäre *nicht* in Ordnung gewesen, weil die Paare aus $\mathbb{N} \times \mathbb{N}$ typischerweise nicht zwei gleiche Komponenten haben.
- Man *interpretiert* einen Ausdruck wie $f(1, 5)$ oft so, dass die Funktion *zwei* Eingaben – nämlich 1 und 5 – bekommt. Aber formal bekommt sie nur *eine* Eingabe: das Paar $(1, 5)$.
- Der letzte Punkt ist insbesondere für die Frage relevant, ob so eine Funktion injektiv ist. Unser Beispiel ist nämlich *nicht* injektiv, weil auch $f(5, 1) = 26$ gilt und $(1, 5)$ und $(5, 1)$ verschiedene Paare sind.

Mit dieser Sichtweise kann man auch ganz „alltägliche“ Operationen wie Addition oder Multiplikation als Funktionen betrachten – und so sieht man sie in der Mathematik auch: Ist M irgendeine Menge, dann bezeichnet man eine Funktion von

mehrstellige Funktion

¹⁴Analog gibt es natürlich auch dreistellige oder allgemein *mehrstellige* Funktionen, deren Argumente Tripel bzw. n -Tupel sind. Man schreibt dann so etwas wie $f(x, y, z)$.

$M \times M$ nach M als (zweistellige) *innere Verknüpfung* auf M . In diesem Sinne ist beispielsweise die Funktion, die (x, y) auf $x + y$ abbildet, eine innere Verknüpfung auf $\mathbb{N}$ (und auch auf $\mathbb{Z}$, $\mathbb{Q}$ oder $\mathbb{R}$).

innere Verknüpfung

Aufgabe 16.17. Wenn man Multiplikation als innere Verknüpfung auf $\mathbb{R}$ betrachtet, ist diese Funktion dann injektiv?

Aufgabe 16.18. Welche der folgenden Funktionen ist injektiv? Begründen Sie negative Antworten jeweils mit einem konkreten Gegenbeispiel.

$$f_1: \begin{cases} \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N} \\ (m, n) \mapsto m \end{cases} \quad f_2: \begin{cases} \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N} \times \mathbb{N} \\ (m, n) \mapsto (n, m) \end{cases} \quad f_3: \begin{cases} \mathbb{N} \rightarrow \mathbb{N} \times \mathbb{N} \\ n \mapsto (n, n) \end{cases}$$

★Aufgabe 16.19. Ist es möglich, dass eine Funktion von $\mathbb{N} \times \mathbb{N}$ nach $\mathbb{N}$ injektiv ist? Geben Sie entweder ein Beispiel für eine solche Funktion an oder begründen Sie, warum das nicht geht.

Für die folgenden Aufgaben sollten Sie sich an Definitionen durch **Fallunterscheidung** erinnern. Außerdem brauchen wir noch ein Konzept, das insbesondere in der Informatik häufig verwendet wird:

Ist x eine reelle Zahl, so schreibt man $\lfloor x \rfloor$ für die größte *ganze* Zahl, die nicht größer als x ist. Ebenso schreibt man $\lceil x \rceil$ für die kleinste *ganze* Zahl, die nicht kleiner als x ist. Das sind die sogenannten **unteren** und **oberen Gaußklammern**.

Definition

Beispielsweise gilt $\lceil 7/2 \rceil = 3 = \lceil 27/10 \rceil$ und $\lfloor -7/2 \rfloor = -4 = \lfloor -47/10 \rfloor$. Man beachte das Verhalten bei negativen Zahlen!

Aufgabe 16.20. Was wäre, wenn man in der **Definition des Absolutbetrags** die Bedingung $x < 0$ durch $x \leq 0$ ersetzen würde?

Aufgabe 16.21. Nur in einem der beiden folgenden Fälle wird wirklich eine Funktion definiert. Wo liegt beim anderen der Fehler?

$$f_1: \begin{cases} \mathbb{N} \rightarrow \mathbb{N} \\ n \mapsto \begin{cases} \lfloor n/10 \rfloor & n \leq 9 \\ \lfloor (n-4)/10 \rfloor & n \geq 4 \end{cases} \end{cases} \quad f_2: \begin{cases} \mathbb{N} \rightarrow \mathbb{N} \\ n \mapsto \begin{cases} \lfloor n/10 \rfloor & n \leq 9 \\ \lfloor (n+4)/10 \rfloor & n \geq 4 \end{cases} \end{cases}$$

Aufgabe 16.22. Schreiben Sie die Funktion

$$f: \begin{cases} [6] \rightarrow \mathbb{N} \\ n \mapsto \begin{cases} n & n \text{ ist gerade} \\ 6-n & n \text{ ist ungerade} \end{cases} \end{cases}$$

als Menge von Paaren auf. Was ist der Wertebereich von f ?

Wenn man bei der Analogie von Funktionen als „Maschinen“ mit Ein- und Ausgabe bleibt, dann kann man sich auch vorstellen, dass solche Maschinen hintereinandergeschaltet werden: Die Ausgabe der einen Maschine wird zur Eingabe der nächsten. Damit hat man gewissermaßen eine neue Maschine erstellt. Führt man beispielsweise erst die Funktion $n \mapsto 5n$ aus und anschließend die Funktion $n \mapsto n + 2$, so erhält man eine neue Funktion, die $n \mapsto 5n + 2$ ausführt. Das ist das Konzept der *Komposition*, das wir nun formal fixieren.

Definition

Ist Sind R und S Relationen, so definieren wir deren **Komposition** durch

$$R \circ S = \{(x, y) : \exists z (xSz \wedge zRy)\}.$$

Diese technische Definition können Sie notfalls vergessen. Für die Anwendung ist die folgende Konsequenz entscheidend:

Satz 16.2

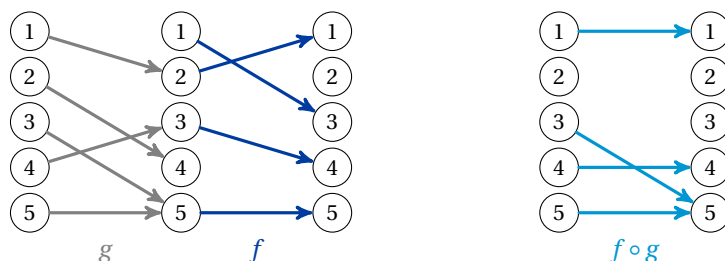
Sind f und g Funktionen, dann ist die Relation $f \circ g$ ebenfalls eine Funktion und für alle x aus dem Definitionsbereich von $f \circ g$ gilt $(f \circ g)(x) = f(g(x))$.

Die inzwischen weltweit übliche Schreibweise $f(x)$ geht übrigens auf den berühmten Schweizer Mathematiker **Leonhard Euler** zurück, der sie 1734 zum ersten Mal verwendete. Weil dabei die Funktion links vom Argument x steht, kommt einem der Ausdruck $f \circ g$ – der zu $f(g(x))$ passt – zunächst „falsch herum“ vor, weil in Europa die rechtsläufige **Schreibrichtung** üblich ist. Gemeint ist nämlich, dass *erst* g und *danach* f ausgeführt wird; die Reihenfolge, in der die Funktionen „drankommen“, verläuft also von rechts nach links.¹⁵

Als übersichtliches Beispiel betrachten wir die Funktionen

$$f = \{(1, 3), (2, 1), (3, 4), (5, 5)\} \quad \text{und} \quad g = \{(1, 2), (2, 4), (3, 5), (4, 3), (5, 5)\}$$

und erhalten die Komposition $f \circ g = \{(1, 1), (3, 5), (4, 4), (5, 5)\}$. Siehe dazu auch die folgende Grafik.



Wenn diese Begriffe neu für Sie sind, dann sollten Sie das Beispiel im Detail nachvollziehen und überprüfen. Dabei sind Ihnen vielleicht zwei Dinge aufgefallen:

¹⁵Aus diesem Grund wird $f \circ g$ auch als „ f nach g “ gelesen; manchmal aber auch als „ f komponiert mit g “, „ f Kuller g “ oder „ f Kringel g “.

- (i) Die Zahl 2 gehört zum Definitionsbereich von g , aber nicht zum Definitionsbereich der Komposition $f \circ g$. Das liegt daran, dass g diese Zahl auf die Zahl 4 abbildet, die aber nicht zum Definitionsbereich von f gehört.
- (ii) Die Zahl 3 gehört zum Wertebereich von f , aber nicht zum Wertebereich der Komposition $f \circ g$. Das liegt daran, dass f die Zahl 1 auf 3 abbildet, 1 jedoch nicht zum Wertebereich von g gehört.

Aufgabe 16.23. Berechnen Sie auch $g \circ f$. Was fällt Ihnen auf?

In der Praxis wird man zumindest die Situation (i) von oben in der Regel vermeiden und nur Kompositionen $f \circ g$ bilden, bei denen der Wertebereich von g eine Teilmenge des Definitionsbereichs von f ist. Grafisch gesprochen soll also jeder Pfeil von g eine „Fortsetzung“ durch einen von f bekommen. Dann haben g und $f \circ g$ denselben Definitionsbereich. So werden wir es im Folgenden auch halten (wenn wir nicht explizit nach seltsamen Gegenbeispielen suchen).

Haben wir für die Funktionen, die komponiert werden, Rechenvorschriften, dann brauchen wir nicht mühsam passende Paare zusammenzusuchen, sondern wir können Satz 16.2 anwenden und eine Rechenvorschrift in die andere „einsetzen“. Haben wir z. B. die durch $f(x) = x^2$ und $g(x) = 3 + 2x$ jeweils auf ganz $\mathbb{R}$ definierten Funktionen, dann ergibt sich die Rechenvorschrift für $f \circ g$ durch

$$(f \circ g)(x) = f(g(x)) = f(3 + 2x) = (3 + 2x)^2 = 9 + 12x + 4x^2.$$

Ganz einfach, wenn man nicht den leider recht häufigen Fehler macht, Klammern zu vergessen. Der Ausdruck für $g(x)$ wird in die Rechenvorschrift für $f(x)$ anstelle von x eingesetzt, aber *als Ganzes*. Vorsichtshalber sollte man ihn dort immer zuerst *mit* Klammern einsetzen. Sollten die sich als unnötig herausstellen, kann man sie immer noch entfernen. Hier wäre die Alternative $3 + 2x^2$ ohne Klammern auf jeden Fall falsch gewesen!

Aufgabe 16.24. Die Funktionen $f(x) = x^2$ und $g(x) = 2^x$ sind auf ganz $\mathbb{R}$ definiert. Geben Sie die Rechenvorschriften für $f \circ g$ und $g \circ f$ an. Begründen Sie mit einem konkreten Gegenbeispiel, dass die beiden Kompositionen nicht gleich sind.

Aufgabe 16.25. Geben Sie zu den Funktionen

$$f: \begin{cases} \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N} \\ (x, y) \mapsto x \end{cases} \quad \text{und} \quad g: \begin{cases} \mathbb{N} \rightarrow \mathbb{N} \times \mathbb{N} \\ x \mapsto (x, x) \end{cases}$$

die Kompositionen $f \circ g$ und $g \circ f$ in derselben ausführlichen Schreibweise wie in der Aufgabenstellung an.

Aufgabe 16.26. Fallen Ihnen zwei Funktionen f und g mit $f \circ g = g \circ f$ ein? (Gemeint sind wirklich *zwei* Funktionen, d. h., es soll $f \neq g$ gelten. Sonst wäre die Lösung „trivial“, wie eine Mathematikerin sagen würde.)

Mithilfe der Komposition können wir Aussagen über die Rolle der Identität und der Umkehrfunktion formulieren, die leicht einzusehen sind, wenn man sich einfach die jeweiligen Definitionen vergegenwärtigt:

Satz 16.3

Sei f eine Funktion mit Definitionsbereich A und Wertebereich B . Dann gilt:

- (i) $\text{id}_B \circ f = f = f \circ \text{id}_A$
- (ii) $f \circ f^{-1} = \text{id}_B$ und $f^{-1} \circ f = \text{id}_A$, falls f injektiv ist.

Man kann das so interpretieren, dass $\circ$ eine Art Multiplikation von Funktionen ist und die Identität die Rolle der Eins spielt, während die Umkehrfunktion quasi der Kehrwert ist.¹⁶ (Das erklärt auch die Schreibweise mit dem Exponenten -1 .)

Es dürfte offensichtlich sein, dass $f \circ g$ injektiv ist, wenn f und g beide injektiv sind. (Machen Sie sich das ggf. mit Pfeildiagrammen klar.) Ist mindestens eine der beiden Funktionen nicht injektiv, dann ist ihre Komposition das im Allgemeinen auch nicht mehr. Nehmen Sie als Beispiel $f(x) = x$ und $g(x) = x^2$ auf den reellen Zahlen: Keine der Funktionen $f \circ g$, $g \circ f$ und $g \circ g$ ist injektiv.

Aufgabe 16.27. Versuchen Sie Funktionen f und g von $\mathbb{N}$ nach $\mathbb{N}$ zu finden, für die eine der folgenden Aussagen gilt.

- (i) f ist nicht injektiv, aber $f \circ g$ ist injektiv.
- (ii) g ist nicht injektiv, aber $f \circ g$ ist injektiv.

Sie werden feststellen, dass das für eine der beiden Aussagen möglich ist und für die andere nicht. Können Sie begründen, warum das so ist?

★Aufgabe 16.28. Finden Sie Funktionen f und g , so dass f und g beide nicht injektiv sind, $f \circ g$ aber injektiv ist. Wieso steht das nicht im Widerspruch zu Aufgabe 16.27? [Tipp: Arbeiten Sie mit ganz einfachen Funktionen, die nur aus wenigen Paaren bestehen.]

Aufgabe 16.29. Welche der vier Funktionen aus Aufgabe 16.25 sind injektiv? Denken Sie wie üblich daran, dass für negative Antworten konkrete Gegenbeispiele am überzeugendsten sind.

Aufgabe 16.30. Welche der vier Funktionen

$$f: \begin{cases} \mathbb{Z} \times \mathbb{Z} \rightarrow \mathbb{Z} \\ (x, y) \mapsto y + 1 \end{cases}, \quad g: \begin{cases} \mathbb{Z} \rightarrow \mathbb{Z} \times \mathbb{Z} \\ x \mapsto (x - 1, x + 1) \end{cases},$$

$f \circ g$ und $g \circ f$ sind injektiv?

Aufgabe 16.31. Welche der vier Funktionen

$$f: \begin{cases} \mathbb{Z} \times \mathbb{Z} \rightarrow \mathbb{Z} \times \mathbb{Z} \\ (x, y) \mapsto (x + y, x - y) \end{cases}, \quad g: \begin{cases} \mathbb{Z} \times \mathbb{Z} \rightarrow \mathbb{Z} \times \mathbb{Z} \\ (x, y) \mapsto (y - 1, x + 1) \end{cases},$$

$f \circ g$ und $g \circ g$ sind injektiv?

¹⁶Natürlich ist die Analogie nicht vollständig, weil es hier viele verschiedene „Einsen“ gibt und weil die Komposition anders als die Multiplikation reeller Zahlen nicht kommutativ ist – siehe Aufgabe 16.24..

Bemerkungen

Vergessen Sie nicht die ungewohnte Reihenfolge bei der Komposition von Funktionen: $f \circ g$ bedeutet, dass erst g , dann f ausgeführt wird.

Durch die Forderung der Rechtseindeutigkeit wird im Nachhinein vielleicht noch einmal klar, warum man bei der Definition der **Wurzel** $\sqrt[n]{x}$ darauf besteht, dass immer die nichtnegative Lösung y von $y^n = x$ gemeint ist. Ansonsten wäre beispielsweise $x \mapsto \sqrt{x}$ keine Funktion, weil die „Maschine“ nicht „wüsste“, ob sie bei der Eingabe 25 die Ausgabe 5 oder die Ausgabe -5 wählen sollte.

Obwohl das in der Schule vielleicht nicht vorkam, können sowohl die Eingaben als auch die Ausgaben von Funktionen andere Objekte als Zahlen sein. Die Funktion g aus Aufgabe 16.12 erwartet beispielsweise als Eingabe eine Menge. Ebenso erhält man durch $n \mapsto [n]$ eine Funktion mit dem Definitionsbereich $\mathbb{N}$, deren Ausgaben Mengen sind. Wenn man Informatik studiert, sollte das eigentlich nicht überraschend sein, weil in den meisten gängigen Programmiersprachen, Funktionen statt mit Zahlen auch mit **Strings**, **Arrays** oder anderen Datentypen umgehen können.

Relationen bilden übrigens die mathematische Grundlage für **relationale Datenbanken**.

★ Funktionen im Computer

In PYTHON sind Funktionen die Objekte, die man mit `def` oder `lambda` erzeugt. Wie oben erwähnt, verhalten die sich nicht immer genauso wie mathematische Funktionen, aber es gibt trotzdem viele Analogien. Insbesondere sind in PYTHON Funktionen ganz „normale“ Objekte, die man beispielsweise wie Zahlen in Datenstrukturen speichern oder sogar als Argumente oder Rückgabewerte anderer Funktionen verwenden kann. Hier ein Beispiel:

```
def square(f):
    return lambda x: f(x)*f(x)

g = lambda x: x+1
h = lambda x: x*x
sh = square(h)

square(g)(3), square(h)(3), sh(2)
```

Die Kompositionen zweier PYTHON-Funktionen lässt sich recht einfach realisieren:

```
def comp(f, g):
    return lambda x: f(g(x))

f, g = lambda x: 3 * x, lambda x: x + 2
comp(f,g)(42)
```

Dass man wie in (16.4) das Argument einer Funktion vorab bereits in seine Einzelteile (im Beispiel sind das m und n) „zerlegt“, nennt man in der Informatik *Destrukturierung*. In PYTHON würden Sie die besagte Funktion typischerweise sicher so schreiben:

```
def f(m,n):  
    return m**2+n**2
```

Dann würden Sie sie in der Form $f(1,5)$ aufrufen. Man könnte das aber auch näher am „Original“ so schreiben:

```
def f(p):  
    (m,n) = p  
    return m**2+n**2
```

Der Aufruf wäre nun $f((1,5))$, d. h., f wird mit einem *Paar* als Argument aufgerufen und das Destrukturieren findet in der ersten Zeile der Funktion statt.¹⁷ In Programmiersprachen wie HASKELL, in denen mit *Pattern Matching* gearbeitet wird, wird diese Vorgehensweise besonders deutlich.

Die Gaußklammern erhält man folgendermaßen:

```
from math import floor, ceil  
floor(-7/2), ceil(-47/10)
```

¹⁷Die Klammern sind in diesem Fall nicht nötig und stehen dort aus didaktischen Gründen.



Geometrie

Die Geometrie ist – zusammen mit Algebra und Analysis – eines der Hauptgebiete der Mathematik.¹ In gewissem Sinne steht sie auch ganz am Anfang der Geschichte der Mathematik, weil die Begründer der wissenschaftlichen Mathematik, die Griechen, hauptsächlich geometrisch dachten. Wir werden uns allerdings nicht aus nostalgischen oder historischen Gründen mit der Geometrie beschäftigen. Vielmehr tauchen geometrische Fragen ganz natürlich in verschiedenen Zusammenhängen auf und außerdem spielt die Geometrie eine zentrale Rolle in vielen Anwendungen, nicht zuletzt in der Computergrafik.

17. Elementargeometrie der Ebene

Zunächst wird es darum gehen, die Geometrie zu wiederholen, die typischerweise an der Schule vermittelt wird. Eigentlich sollten Sie das also alles wissen, aber wahrscheinlich haben Sie einiges vergessen. Außerdem kann es sicherlich nicht schaden, es noch einmal im Zusammenhang zu sehen und Begriffe zu präzisieren. Wir werden allerdings nicht so präzise sein, wie man das in einem Studium der Mathematik wäre. Insbesondere werden wir die Geometrie nicht **axiomatisch** aufbauen, sondern bei grundlegenden Begriffen wie *Punkt* oder *Gerade* auf anschaulichen Vorstellungen aufbauen.

Halten wir jedoch am Anfang zumindest fest, dass mathematische Geometrie immer idealisiert und nie über reale Objekte spricht. Die Geometrie in diesem Abschnitt spielt sich in einer Ebene ab, die unendlich groß ist – der sogenannten *euklidischen Ebene*. Es gibt also nirgends eine Grenze oder ein „Ende der Welt“. Obwohl es momentan noch keine große Rolle spielen wird, ist es bereits jetzt sinnvoll, den Zusammenhang zu **Mengen** herzustellen. In diesem Sinne ist die euklidische Ebene eine Menge, die aus unendlich vielen *Punkten* besteht. Punkte beschreiben Orte in der Ebene und haben keine Ausdehnung. Alle geometrischen Objekte, mit denen wir uns im Folgenden beschäftigen werden, zum Beispiel Dreiecke oder

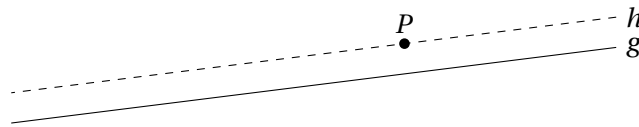
euklidische Ebene

Punkt

¹Wobei das nicht so gemeint ist, dass diese Gebiete streng voneinander abgegrenzt werden können. Auch gibt es natürlich noch viele andere Bereiche der Mathematik, die man keinem der drei direkt zuordnen kann. Es ist aber eine grobe Unterteilung, die als erste Annäherung ganz gut funktioniert.

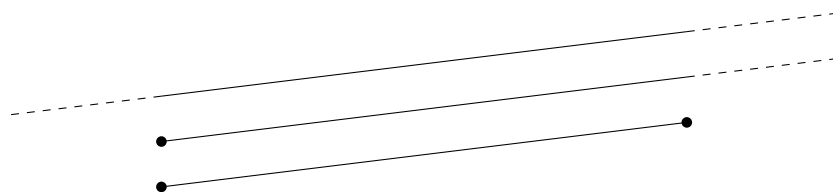
Kreise, sind Mengen von Punkten. Wenn wir z. B. sagen, dass ein Punkt *auf* einem Kreis liegt, dann ist damit gemeint, dass der Punkt ein *Element* der Menge der Punkte ist, die den Kreis bilden. Wenn wir sagen, dass zwei Geraden sich *schneiden*, dann soll das bedeuten, dass der *Durchschnitt* der beiden Punktmengen, die die beiden Geraden bilden, nicht leer ist. Und so weiter.

- Gerade Die grundlegendsten Objekte neben Punkten sind *Geraden*. Geraden sind unendlich lang, haben aber keine Breite. Man kann sich jede Gerade wie eine Kopie der *reellen Zahlengeraden* vorstellen, wobei die Zahlen den Punkten auf der Geraden entsprechen. Insbesondere haben wir die folgenden wesentlichen Eigenschaften:
- (i) Zwischen je zwei *verschiedenen* Punkten auf einer Geraden liegen unendlich viele weitere Punkte.
 - (ii) Zu je zwei *verschiedenen* Punkten der Ebene gibt es *genau eine* Gerade durch diese beiden Punkte. (In Mengen „übersetzt“ heißt das, dass diese Gerade die beiden vorgegebenen Punkte als Elemente enthält.)
 - (iii) Zwei *verschiedene* Geraden in der Ebene schneiden sich entweder in genau einem Punkt oder sie haben keinen Punkt gemeinsam. Anders ausgedrückt: Wenn sie mehr als einen Punkt gemeinsam haben, sind sie identisch. Das folgt aus (ii).
- parallel (iv) Zwei Geraden werden als *parallel* bezeichnet, wenn sie entweder identisch sind oder keinen Punkt gemeinsam haben. Zu jeder Geraden g und jedem Punkt P , der *nicht* auf g liegt, gibt es *genau eine* Gerade h , die parallel zu g ist und P enthält.



- Halbgerade / Strahl Ein Punkt P auf einer Geraden teilt diese in zwei Teile auf, die man *Halbgeraden* oder *Strahlen* nennt. Wenn man sich die Gerade wie die Menge $\mathbb{R}$ vorstellt und den Punkt P wie eine reelle Zahl a , dann ist eine Halbgerade ein *Intervall* der Form $(-\infty, a]$ oder $[a, \infty)$. Ein Stück einer Geraden zwischen zwei *verschiedenen* Punkten, die auf ihr liegen, nennt man eine *Strecke*. Das entspricht dann einem Intervall der Form $[a, b]$.
- Strecke

In der folgenden Skizze werden eine Gerade, ein Strahl und eine Strecke untereinander dargestellt. Die gestrichelten Teile sollen dabei andeuten, dass das jeweilige Objekt sich in dieser Richtung unendlich weit ausdehnt.

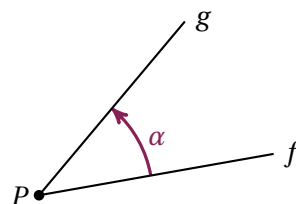


Zwei Anmerkungen dazu:

- (i) Skizzen wie diese können nie die idealisierte mathematische „Wirklichkeit“ darstellen, da in dieser Punkte keine Ausdehnung und Geraden keine Breite haben: Wir könnten sie gar nicht sehen. Die angeblichen Punkte in der Skizze sind beispielsweise eigentlich kleine Kreisscheiben und durch zwei solche „Punkte“ gibt es natürlich mehr als eine Gerade. Diese Beobachtung ist nicht nur Pedanterie, sondern unter anderem relevant für die Computergrafik, in der man zwar die mathematische Geometrie einsetzt, aber immer im Kopf haben muss, dass das, was auf dem Bildschirm angezeigt wird, nur eine Approximation ist.
- (ii) Nach der obigen Einführung gehören die Anfangs- und Endpunkte zu einer Strecke bzw. einem Strahl dazu. Aber wie man bei Intervallen zwischen $[a, b]$, $[a, b)$ und (a, b) unterschieden kann, kann man auch in der Geometrie Strahlen und Strecken *ohne* Randpunkte betrachten. Im Folgenden werden wir das jedoch in der Regel nicht benötigen.

Auf den folgenden Seiten werden wir Konzepte wie die *Länge* von Strecken oder später den *Flächeninhalt* einfacher Figuren nicht weiter hinterfragen. Die anschauliche Vorstellung davon, was damit gemeint ist, wird für unsere Zwecke ausreichen.² Etwas genauer betrachten werden wir allerdings ein Konzept, das auf den ersten Blick ebenfalls einfach erscheint, erfahrungsgemäß aber durchaus Probleme bereiten kann: das des *Winkels*. (Vielleicht überlegen Sie erst einmal selbst, wie Sie ohne Zeichnungen und Gesten erklären würden, was genau ein Winkel ist!) Wir werden die folgende (anschauliche) Definition verwenden: Sind f und g zwei Strahlen mit demselben Anfangspunkt P , dann gibt der *gerichtete* bzw. *orientierte Winkel* (engl. *signed angle*) zwischen f und g die Drehung um den Punkt P an, die f mit g zur Deckung bringt. P nennt man in diesem Zusammenhang den *Scheitelpunkt* (engl. *vertex*) des Winkel und f und g dessen *Schenkel* (engl. *sides*).

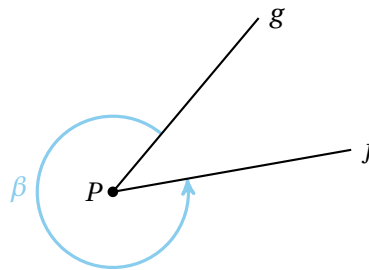
gerichteter Winkel



Wenn nicht explizit etwas anderes gesagt wird, ist bei Drehungen immer die *mathematisch positive* Drehrichtung gemeint, die *gegen* den Uhrzeigersinn verläuft.³ Beachten Sie, dass nach unserer Definition der Winkel zwischen g und f ein anderer ist als der zwischen f und g !

²In der höheren Mathematik kann die genaue Klärung dieser Begriffe ziemlich kompliziert werden. Das werden wir jedoch nicht brauchen.

³Eselsbrücke: Mathematik gibt es schon wesentlich länger als Uhren, also laufen Uhren eigentlich gegen den mathematischen Drehsinn ...



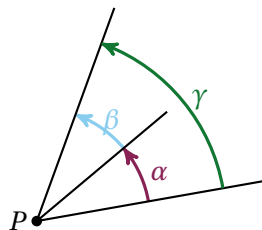
Grad

rechter Winkel

Um mit Winkeln arbeiten zu können, ordnet man ihnen Zahlen zu. Wir werden zunächst das allgemein gebräuchliche *Gradmaß* verwenden, das dem Winkel, der einer kompletten Umdrehung entspricht, den Wert 360 Grad – in Zeichen: 360° – zuordnet.⁴ Der Hälfte einer solchen Umdrehung entsprechen dann 180° und so weiter. Die folgende Tabelle zeigt ein paar wesentliche Winkel. Den Winkel, der 90° entspricht, nennt man bekanntlich einen *rechten Winkel*.

90°	180°	270°	360°

An diesen Beispielen sieht man bereits, dass man sinnvoll mit Winkeln rechnen kann: 90 Grad entsprechen z. B. einem Viertel von 360 Grad, weil man nur ein Viertel der vollen Drehung durchführt. Man kann also Winkel mit Zahlen multiplizieren.⁵ Ebenso ergibt es Sinn, Winkel zu addieren. Die folgende Skizze verdeutlicht hoffentlich die Beziehung $\gamma = \alpha + \beta$.



Und natürlich kann man mit derselben Berechtigung auch subtrahieren: $\beta = \gamma - \alpha$. Wie im [Kapitel über Arithmetik](#) kann man $\gamma - \alpha$ auch als $\gamma + (-\alpha)$ interpretieren und den Winkel $-\alpha$ als den zu α inversen Winkel betrachten. Anschaulich steht $-\alpha$ für eine Drehung *im Uhrzeigersinn*, wenn α eine gegen den Uhrzeigersinn ist – und umgekehrt.

Anders als beim Rechnen mit Zahlen gibt es bei Winkeln jedoch die Besonderheit, dass eine Drehung um 360° wieder zur Ausgangslage zurückführt und damit einer „Drehung“ um 0° entspricht. Und das gilt auch für mehrfaches Drehen um 360°

⁴Aus mathematischer Sicht ist diese Zahl völlig willkürlich gewählt. Man hätte auch 42, 100 oder 2025 nehmen können. Dass es 360 ist, liegt daran, dass bereits die [Babylonier](#) diese Einteilung verwendet haben und sie von den Astronomen späterer Kulturen übernommen wurde.

⁵Wo hingegen es keine passende Interpretation für das Produkt zweier Winkel gibt.

bzw. um -360° . Mit anderen Worten: Unterscheiden sich die Gradmaße zweier Winkel um ein ganzzahliges Vielfaches von 360° , dann sind die Winkel identisch.

Aufgabe 17.1. Geben Sie für die folgenden Winkel jeweils ein Gradmaß an, das zwischen 0° und 360° liegt: $\alpha = 740^\circ$, $\beta = -20^\circ$, $\gamma = -400^\circ$, $\delta = 400.3^\circ$.

Für die Richtung, die in der obigen Definition von Winkeln vorkommt, muss man angeben, von wo nach wo gedreht werden soll. Häufig benötigt man diese Richtung jedoch gar nicht, sondern nur den Betrag des Winkelmaßes. Man spricht dann vom *ungerichteten* Winkel zwischen g und h und dieser ist identisch mit dem ungerichteten Winkel zwischen h und g . Ungerichtete Winkel sind natürlich niemals negativ. Wenn man zwischen diesen beiden Sorten von Winkeln unterscheiden will, verwendet man in der Regel das Zeichen $\sphericalangle$ für gerichtete und $\sphericalangle$ für ungerichtete Winkel. Das ist aber eine von den vielen Konventionen, an die sich nicht alle halten.

ungerichteter Winkel

$\sphericalangle$ $\sphericalangle$

Es folgen ein paar einfache Eigenschaften von (ungerichteten) Winkeln. Die dort auftauchenden Begriffe müssen Sie sich nicht alle merken. Aber die beschriebenen Zusammenhänge sollten Ihnen klar sein.

Der Winkel zwischen zwei Halbgeraden, die auf derselben Geraden liegen, hat das Winkelmaß 180° . Teilt man solch einen Winkel in zwei (nicht notwendig gleich große) Teile, so spricht man von *Ergänzungswinkeln*. α und β in der folgenden Skizze sind also Ergänzungswinkel und ihre Summe beträgt 180° .

Ergänzungswinkel



Sind zwei Ergänzungswinkel gleich groß, so handelt es sich um rechte Winkel. In Deutschland und den meisten anderen europäischen Ländern symbolisiert man einen rechten Winkel häufig durch einen Viertelkreis mit einem kleinen Punkt drin. Im anglophonen Raum verwendet man stattdessen ein kleines Quadrat.



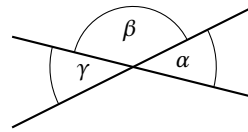
Winkel, die kleiner als rechte Winkel sind, nennt man *spitz*. Winkel, die größer als rechte und kleiner als 180 Grad sind, heißen *stumpf*. (Die entsprechenden englischen Begriffe sind *acute* und *obtuse*.)

spitz
stumpf

Wenn man zwei Geraden schneidet, erhält man insgesamt vier Winkel. Liegen zwei von denen aneinander, so nennt man sie *Nebenwinkel*. Sie sind dann offenbar Ergänzungswinkel. Liegen zwei der vier nicht aneinander, so liegen sie sich gegenüber und werden *Scheitel-* oder *Gegenwinkel* genannt. Gegenwinkel sind immer gleich groß. In der folgenden Skizze sind z. B. α und β Nebenwinkel und α und γ Gegenwinkel, also gilt $\alpha + \beta = 180^\circ$ und $\alpha = \gamma$.

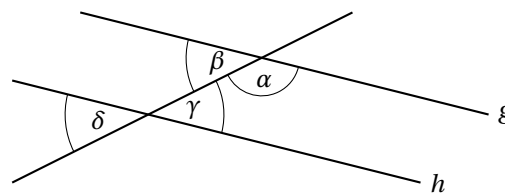
Nebenwinkel

Scheitelwinkel



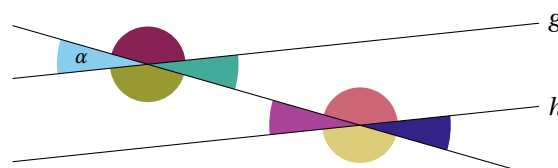
Hat man es mit zwei parallelen Geraden und einer dritten, die diese beiden schneidet, zu tun, so ergeben sich auch wieder spezielle Winkelbeziehungen.

- | | |
|---------------|---|
| Stufenwinkel | <ul style="list-style-type: none"> • Zwei Winkel auf derselben Seite der schneidenden Gerade und auf entsprechenden Seiten der parallelen Geraden heißen <i>Stufenwinkel</i>. In der folgenden Skizze sind das z. B. β und δ. Stufenwinkel sind immer gleich groß, d. h., hier gilt $\beta = \delta$. |
| Wechselwinkel | <ul style="list-style-type: none"> • Liegen zwei Winkel auf entgegengesetzten Seiten der parallelen Geraden und auf verschiedenen Seiten der schneidenden Gerade, so heißen sie <i>Wechselwinkel</i>. In der Skizze trifft das auf β und γ zu. Wechselwinkel sind ebenfalls immer gleich groß: $\beta = \gamma$. |
| Nachbarwinkel | <ul style="list-style-type: none"> • Schließlich nennt man zwei Winkel auf derselben Seite der schneidenden Gerade und auf unterschiedlichen Seiten der parallelen Geraden <i>Nachbarwinkel</i>. In der Skizze unten sind das etwa α und γ. Nachbarwinkel ergänzen sich immer zu 180 Grad: $\alpha + \gamma = 180^\circ$. |

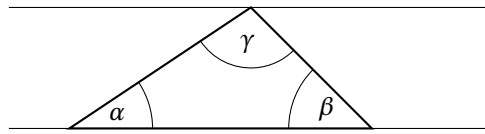


Umgekehrt kann man aus der Gleichheit entsprechender Winkel auch auf die Parallelität schließen. Sind in der obigen Skizze beispielsweise β und δ identisch, dann sind g und h parallel.

Aufgabe 17.2. In der folgenden Skizze sollen g und h parallel sein. Der hellblaue Winkel α links beträgt 22° . Geben Sie die Maße der anderen sieben Winkel an.



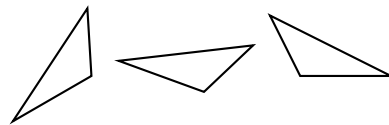
- | | |
|---------|---|
| Dreieck | <p>Die einfachste geometrische Figur ist das <i>Dreieck</i>. Die „Zutaten“ dafür sind drei Punkte, die <i>nicht</i> auf einer gemeinsamen Geraden liegen, sowie die Strecken, die jeweils zwei dieser Punkte verbinden. Das sind die <i>Ecken</i> bzw. <i>Seiten</i> des Dreiecks. Nach den obigen Überlegungen folgt, dass die Summe der Innenwinkel in <i>jedem</i> Dreieck 180° beträgt.</p> |
|---------|---|



Aufgabe 17.3. Begründen Sie mithilfe der obigen Skizze, dass $\alpha + \beta + \gamma = 180^\circ$ gilt.

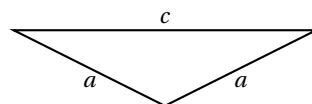
Eine ganz wesentliche Beobachtung ist, dass sich die Eigenschaften von Dreiecken (wie etwa die Fläche, die Längen der Seiten, die Winkel der Seiten zueinander) nicht ändern, wenn man sie verschiebt, dreht oder spiegelt. (Das ist hoffentlich intuitiv klar.) Man sagt dann, dass die Dreiecke *kongruent*⁶ zueinander sind – so wie diese drei:

kongruent



Man kann das auch umdrehen: wie viele und welche Eigenschaften von zwei Dreiecken müssen identisch sein, damit die beiden kongruent sind? Diese Frage wird von den *Kongruenzsätzen* beantwortet, die Sie vielleicht unter ihren abkürzenden Namen wie SWS oder WSW in Erinnerung haben. SSS besagt z. B., dass zwei Dreiecke kongruent sind, wenn sie in allen drei Seitenlängen übereinstimmen.

Aufgabe 17.4. Schauen Sie sich (z. B. auf Wikipedia) an, welche Kongruenzsätze es gibt. Versuchen Sie für mindestens einen dieser Sätze, das entsprechende Dreieck geometrisch zu konstruieren. (Für SSS würde das z. B. bedeuten, dass Sie sich drei Längen vorgeben und dann mit Zirkel und Lineal ein Dreieck konstruieren, dessen Seiten exakt diese drei Längen haben.) An der Konstruktion sollte man erkennen können, dass es nur eine Möglichkeit gibt, so ein Dreieck zu erhalten.⁷



Aufgabe 17.5. Ein *gleichschenkliges Dreieck* ist eines, bei dem zwei Seiten gleich lang sind. Begründen Sie mithilfe des Kongruenzsatzes SSS, dass in solchen Dreiecken auch die beiden diesen Seiten gegenüberliegenden Winkel gleich groß sind.

gleichschenkliges Dreieck

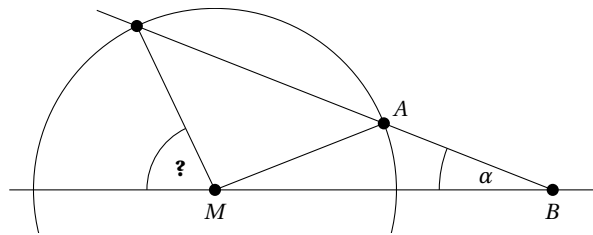
Aufgabe 17.6. Ein *gleichseitiges Dreieck* ist eines, bei dem alle drei Seiten gleich lang sind. Anschaulich ist klar, dass dann auch alle drei Winkel gleich sind. Aber können Sie das auch begründen? Und wie groß sind die Winkel?

gleichseitiges Dreieck

Aufgabe 17.7. In der folgenden Skizze ist M der Mittelpunkt des Kreises. Die Strecken $[AM]$ und $[AB]$ sind gleich lang. Geben Sie den durch ein Fragezeichen markierten Winkel in Abhängigkeit von α an.

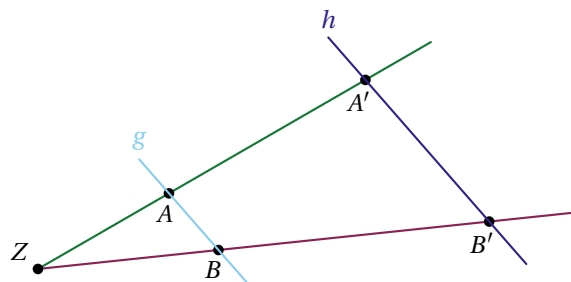
⁶Das bedeutet so viel wie *deckungsgleich* oder *gleichförmig*.

⁷Genauer sollte man sagen: Es gibt nur eine Möglichkeit *modulo Kongruenz*. Damit ist gemeint, dass Sie sich die Lage einer Seite und die Orientierung der Seiten zueinander zwar aussuchen können, der Rest sich dann jedoch zwangsläufig ergibt.



Strahlensatz

Noch wichtiger als die Kongruenz ist jedoch das Konzept der *Ähnlichkeit*, das sich aus dem sogenannten *Strahlensatz* ergibt, den die folgenden Skizze illustriert.



[XY]

[XY]

Von einem Punkt Z gehen zwei Strahlen aus, die von zwei parallelen Geraden g und h (die nicht durch Z gehen) geschnitten werden. Dadurch ergeben sich auf dem einen Strahl die Strecken $[ZA]$, $[ZA']$ und $[AA']$, denen die Strecken $[ZB]$, $[ZB']$ und $[BB']$ auf dem anderen Strahl entsprechen.⁸ Natürlich werden einander entsprechende Strecken nicht unbedingt dieselbe Länge haben. Es gilt also im Allgemeinen *nicht* so etwas wie $|ZA| = |ZB|$ oder $|AA'| = |BB'|$. Der Strahlensatz besagt jedoch, dass die *Verhältnisse* der Längen zueinander identisch sind. Mit den Bezeichnungen aus der Skizze gilt also

$$\frac{|ZA|}{|ZA'|} = \frac{|ZB|}{|ZB'|}, \quad \frac{|ZA|}{|AA'|} = \frac{|ZB|}{|BB'|} \quad \text{und} \quad \frac{|ZA'|}{|AA'|} = \frac{|ZB'|}{|BB'|}. \quad (17.1)$$

Zusätzlich entspricht das erste dieser drei Verhältnisse auch noch dem der Abschnitte auf den Parallelen, also

$$\frac{|ZA|}{|ZA'|} = \frac{|ZB|}{|ZB'|} = \frac{|AB|}{|A'B'|}.$$

Außerdem gilt auch die Umkehrung: Aus jeder der drei Gleichungen in (17.1) folgt, dass die Gerade durch A und B parallel zu der durch A' und B' ist.⁹

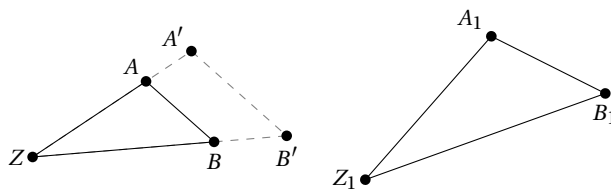
Aufgabe 17.8. Vielleicht ist es keine schlechte Idee, sich noch einmal zu vergewissern, dass man das Konzept des *Verhältnisses* (eigentlich nur ein anderes Wort für *Quotient*) korrekt verstanden hat. Dafür betrachten wir die vier Zahlen $x_1 = 2.4$, $x_2 = 3.2$, $x_3 = 3.6$ und $x_4 = 4.8$. Welche Paare dieser Zahlen haben dasselbe Verhältnis zueinander?

⁸Wir verwenden hier die naheliegende Schreibweise $[XY]$ für die Strecke, die Endpunkte X und Y hat. Und direkt danach die Schreibweise $|XY|$ für die Länge so einer Strecke.

⁹Die Voraussetzung dabei ist, dass ein Strahl von Z durch A und A' verläuft und einer von Z durch B und B' . Wird die zweite oder dritte Gleichung verwendet, muss zusätzlich gefordert werden, dass A wie in der Skizze zwischen Z und A' liegt und B zwischen Z und B' .

★**Aufgabe 17.9.** Machen Sie sich klar, dass aus jeder der drei Aussagen in (17.1) die anderen beiden folgen.

Mit dem Strahlensatz können wir nun folgern, dass bei der maßstabsgetreuen Vergrößerung oder Verkleinerung eines Dreiecks alle Winkel erhalten bleiben. Genauer formuliert: Wenn wir zwei Dreiecke mit den Ecken Z , A und B bzw. Z_1 , A_1 und B_1 haben, so dass die Verhältnisse der Streckenlängen $|ZA|$ zu $|Z_1A_1|$, $|AB|$ zu $|A_1B_1|$ und $|BZ|$ zu $|B_1Z_1|$ alle gleich sind, dann entspricht der Winkel im ersten Dreieck im Eckpunkt Z dem des zweiten im Eckpunkt Z_1 und das gilt auch für die Eckenpaare A und A_1 bzw. B und B_1 .



Um das zu begründen, zeichnet man wie in der obigen Skizze das erste Dreieck und verlängert¹⁰ die Seiten $[ZA]$ und $[ZB]$ bis zu Punkten A' und B' , so dass $[ZA']$ dieselbe Länge hat wie $[Z_1A_1]$ und $[ZB']$ dieselbe wie $[Z_1B_1]$. Aus der Umkehrung des Strahlensatzes folgt wegen der Gleichheit der Streckenverhältnisse, dass die beiden Geraden durch A und B bzw. durch A' und B' parallel sind. Daher stehen auch die Längen $|AB|$ und $|A'B'|$ im selben Verhältnis zueinander wie z. B. $|ZA|$ zu $|ZA'|$. Also hat das Dreieck mit den Ecken Z , A' und B' dieselben Seitenlängen wie das mit den Ecken Z_1 , A_1 und B_1 . Nach dem Kongruenzsatz SSS haben diese beiden Dreiecke dann auch dieselben Winkel. Und das sind, wie man am linken Teil der Skizze erkennt, wiederum dieselben Winkel wie im Dreieck mit den Ecken Z , A und B .

So eine Begründung, die auf den ersten Blick durchaus verwirrend sein kann, müssen Sie nicht eigenständig reproduzieren können, aber Sie sollten sie zumindest verstehen. Mit solchen und ähnlichen Argumenten kommt jedenfalls auf das folgende wichtige Resultat:

Für zwei Dreiecke ABC und $A'B'C'$ sind die folgenden Aussagen äquivalent:

- (i) Die entsprechenden Streckenverhältnisse – wie z. B. $|AB|$ zu $|A'B'|$ – sind alle gleich.
- (ii) Die entsprechenden Winkel – wie z. B. der bei A und der bei A' – sind alle gleich.

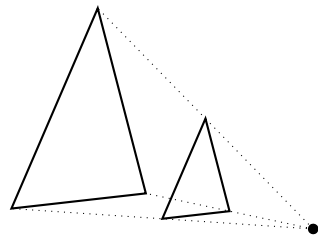
Satz 17.1

In so einem Fall nennt man die beiden Dreiecke *ähnlich*. Wie oben schon erwähnt kann man das auch so interpretieren, dass eines der beiden eine maßstabsgetreue „Kopie“ des anderen ist. Oder dass es wie in der folgenden Skizze *zentrische*

ähnlich

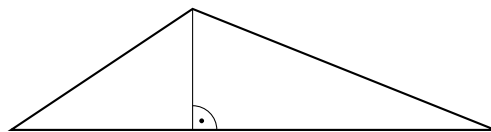
¹⁰O. B. d. A. gehen wir davon aus, dass das zweite Dreieck größer als das erste ist. Anderenfalls vertauschen wir einfach die Rollen.

Streckung hervorgeht (nachdem es vorher ggf. noch verschoben, gedreht und gespiegelt wurde, damit es „passend“ liegt).



Aufgabe 17.10. Warum gibt es keinen Kongruenzsatz WWW?

Nun machen wir es uns noch einfacher und betrachten nur noch *rechtwinklige Dreiecke*, also solche, bei denen ein Winkel ein rechter ist.¹¹ Man kann jedes beliebige Dreieck in zwei rechtwinklige zerlegen, indem man vom größten Winkel aus ein Lot auf die gegenüberliegende Seite fällt.¹² Häufig kann man ein allgemeines Dreieck auf diese Art zerlegen und dann auf die beiden Teile Aussagen über rechtwinklige Dreiecke anwenden.



Hypotenuse
Kathete

In einem rechtwinkligen Dreieck ist der rechte Winkel der größte und dementsprechend die ihm gegenüberliegende Seite die längste. Man nennt sie *Hypotenuse*. Die beiden anderen Seiten heißen *Katheten*.

Die zentrale Aussage über rechtwinklige Dreiecke ist die folgende:¹³

Satz 17.2

Satz des Pythagoras

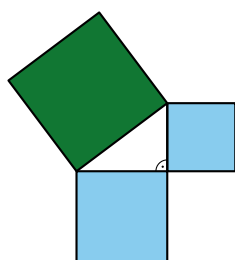
In jedem rechtwinkligen Dreieck entspricht die Summe der Flächen von zwei Quadraten, deren Seitenlängen die der Katheten sind, der Fläche des Quadrates, dessen Seitenlänge die der Hypotenuse ist.

Ein Beispiel dafür zeigt die folgende Skizze. Der Satz besagt, dass die Fläche der beiden kleinen blauen Quadrate zusammen die Fläche des großen grünen ergibt.

¹¹Weil die Winkelsumme 180° beträgt, kann maximal ein Winkel ein rechter Winkel sein.

¹²Man kann auch einen anderen Winkel wählen, aber dann liegt der Fußpunkt des Lots ggf. außerhalb des Dreiecks. In so einem Fall hätte man das Dreieck quasi als „Differenz“ von zwei rechtwinkligen dargestellt, was aber manchmal auch hilfreich sein kann.

¹³Der mit Abstand bekannteste mathematische Satz. Die Forschung ist sich übrigens nicht sicher, ob *Pythagoras*, über den man insgesamt sehr wenig weiß, diesen Satz bewiesen hat. Mit Sicherheit war der Satz selbst aber schon lange vor seiner Zeit bekannt.

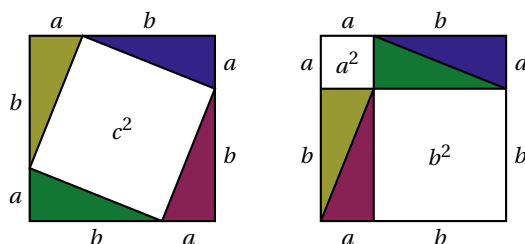


Bezeichnet man die Länge der Hypotenuse mit c und die Längen der Katheten mit a und b , so wird daraus die Formel

$$a^2 + b^2 = c^2, \quad (17.2)$$

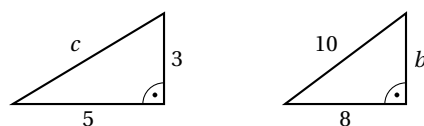
weil die Fläche eines Quadrates sich ergibt, wenn man die Länge einer Seite mit sich selbst multipliziert.¹⁴ Man sollte sich jedoch klarmachen, dass dies *nicht* der Satz des Pythagoras ist, obwohl das häufig die erste Antwort ist, die man bekommt, wenn man nach dem Satz fragt. (17.2) ist eine inhaltsleere Formel (eine *Aussageform* mit drei freien Individuenvariablen), die erst dann Sinn ergibt, wenn man sagt, wofür a , b und c stehen. Der Satz des Pythagoras ist eine geometrische Aussage und keine Formel!

Für den Satz gibt es Hunderte von Beweisen, die mehr oder weniger mathematisch streng sind. Hier einer, der (hoffentlich) ohne Worte auskommt:



An der Grafik sieht man nebenbei, dass zwei rechtwinklige Dreiecke mit den Katheten a und b zusammen ein Rechteck mit den Seiten a und b ergeben. Da die Fläche des Rechtecks ab ist, muss die Fläche so eines Dreiecks $ab/2$ sein.

Der Satz des Pythagoras ermöglicht uns, die Länge einer Seite in einem rechtwinkligen Dreieck zu berechnen, wenn wir die Längen der beiden anderen Seiten kennen. Im linken Teil der folgenden Skizze sind z. B. die Längen der Katheten bekannt. Nach dem Satz gilt $c^2 = 5^2 + 3^2 = 34$ für die Hypotenuse c , also hat c die Länge $\sqrt{34}$.¹⁵



¹⁴Deshalb nennt man das auch *Quadrieren*.

¹⁵Die Gleichung $c^2 = 34$ hat zwei Lösungen, nämlich $\sqrt{34}$ und $-\sqrt{34}$. Da aber die Länge einer Seite gesucht ist, die natürlich nicht negativ sein kann, kommt nur eine der beiden Lösungen in Frage.

Aufgabe 17.11. Ermitteln Sie die Länge von b in dem rechtwinkligen Dreieck im rechten Teil der obigen Skizze. (Beachten Sie, dass die beiden Dreiecke nicht im selben Maßstab gezeichnet wurden.)

Aufgabe 17.12. In der Skizze auf Seite 146 wurde gezeigt, wie man ein beliebiges Dreieck in zwei rechtwinklige Dreiecke zerlegen kann. Begründen Sie, wie daraus die bekannte Formel für die Fläche eines Dreiecks folgt, die die sogenannte *Höhe* des Dreiecks verwendet. (Wie man die Fläche eines rechtwinkligen Dreiecks berechnet, wurde weiter oben beschrieben.)

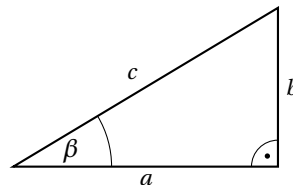
Übrigens gilt auch die folgende Umkehrung des Satzes des Pythagoras:

Lemma 17.3

Wenn für die Seiten a , b und c eines Dreiecks die Beziehung (17.2) gilt, dann ist der Winkel gegenüber der Seite mit der Länge c ein rechter.

Wenn wir in einem rechtwinkligen Dreieck einen der beiden kleinen Winkel kennen, so kennen wir (weil die Winkelsumme 180° betragen muss) auch den anderen. Und damit sind nach den vorherigen Überlegungen auch alle Seitenverhältnisse festgelegt. Es ist daher sinnvoll, diesen nur von einem Winkel abhängenden Quotienten Namen zu geben. Dazu wählt man sich einen Winkel aus (den wir hier β nennen wollen) und nennt die an diesem Winkel anliegende Kathete *Ankathete* und die dem Winkel gegenüberliegende Seite *Gegenkathete*. In der folgenden Grafik ist b die Gegenkathete zu β und a die Ankathete.

An- und Gegenkathete



Sinus / Kosinus

Man definiert nun den *Sinus* von β als das Seitenverhältnis b/c („Gegenkathete durch Hypotenuse“) und den *Kosinus* als a/c ; zusätzlich definiert man noch den *Tangens* von β als b/a . Die Schreibweisen dafür sind $\sin \beta$, $\cos \beta$ und $\tan \beta$. Ist der Winkel ein zusammengesetzter Ausdruck, so benutzt man in der Regel Klammern: $\sin(\alpha + 2\beta)$. Außerdem hat sich eine abkürzende Schreibweise für Potenzen solcher Ausdrücke eingebürgert: man schreibt den Exponenten direkt an den Namen der Funktion. Statt $(\cos \alpha)^3$ schreibt man also z. B. $\cos^3 \alpha$. Wenn Sie daher in einem Fachtext einen Ausdruck wie $\cos \alpha^3$ sehen, ist damit typischerweise $\cos(\alpha^3)$ und *nicht* $(\cos \alpha)^3$ gemeint.¹⁶

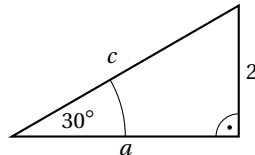
Tangens

Durch einfaches Kürzen folgt sofort, dass der Tangens der Quotient von Sinus und Kosinus ist. Er ist also nicht wirklich „nötig“, sondern könnte als abgeleitete Größe betrachtet werden. Insgesamt kann man sechs „interessante“ Seitenverhältnisse

¹⁶Zur Schreibweise siehe auch die Bemerkungen am Ende des Abschnitts.

(und drei uninteressante wie z. B. b/b) definieren. In der Regel kommt man aber mit den drei genannten aus.¹⁷

Diese sogenannten *trigonometrischen Funktionen* sind deswegen so wichtig, weil man, wenn man in einem rechtwinkligen Dreieck einen Winkel und die Länge einer Seite kennt, mit ihrer Hilfe die Längen der anderen beiden Seiten ermitteln kann. Im Dreieck unten kennen wir z. B. zunächst nur den Winkel 30° sowie die Länge 2 der zugehörigen Gegenkathete.



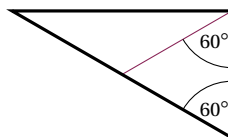
Wenn Sie nun wissen, dass der Sinus von 30° (also das Verhältnis $2/c$) den Wert $1/2$ hat, dann folgt sofort $c = 4$. Wissen Sie außerdem, dass der Kosinus von 30° (das Verhältnis a/c) den Wert $\sqrt{3}/2$ hat, so erhalten Sie $a = 2\sqrt{3}$.

Aufgabe 17.13. Zeichnen Sie mit dem Geodreieck ein rechtwinkliges Dreieck, bei dem die Hypotenuse die Länge 7 und ein Winkel das Maß $\beta = 20^\circ$ hat. Berechnen Sie die Längen der beiden Katheten, wobei Sie für den Sinus bzw. den Kosinus von β die Näherungswerte 0.342 bzw. 0.940 verwenden können. Vergleichen Sie die berechneten Längen mit Ihrer Zeichnung.

Aufgabe 17.14. Zeichnen Sie mit dem Geodreieck ein rechtwinkliges Dreieck, bei dem ein Winkel das Maß $\beta = 40^\circ$ und seine Ankathete die Länge 4 hat. Berechnen Sie die Längen der beiden anderen Seiten, wobei Sie für den Tangens von β die Näherung 0.839 verwenden können. Verwenden Sie *keine* anderen trigonometrischen Funktionen. Vergleichen Sie die berechneten Längen mit Ihrer Zeichnung.

Aufgabe 17.15. Überlegen Sie sich (ohne einen Computer zur Hilfe zu nehmen), welchen Wert der Tangens von 45° haben muss. Machen Sie sich dazu ggf. eine Skizze. Wie groß ist der dritte Winkel in so einem rechtwinkligen Dreieck? Und wie bekommt man dann den Sinus und den Kosinus dieses Winkels? Benutzen Sie dafür den Satz des Pythagoras.

Aufgabe 17.16. Ermitteln Sie (ohne Computer) Sinus, Kosinus und Tangens von 60° . Benutzen Sie dafür die rot eingezeichnete Hilfslinie in der folgenden Skizze.



Aufgabe 17.17. Und wie ist es mit den Werten für 30° ? Das sollte mithilfe der vorherigen Aufgabe jetzt ganz leicht sein.

¹⁷Fall es Sie interessiert: Die anderen drei sind *Sekans*, *Kosekans* und *Kotangens*, die sich als Kehrwerte der drei bereits definierten Verhältnisse darstellen lassen.

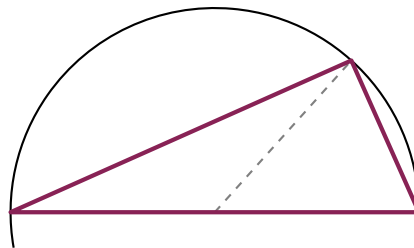
Da rechtwinklige Dreiecke so praktisch sind, ist der folgende Satz, der auch schon seit der griechischen Antike bekannt ist, oft hilfreich:

Satz 17.4

Satz des Thales

Ein Dreieck, das dadurch entsteht, dass man den Durchmesser eines Kreises mit einem weiteren Punkt auf dem Kreis verbindet, ist immer rechtwinklig. Umgekehrt verläuft der **Umkreis** eines rechtwinkligen Dreiecks (also der Kreis, der durch die drei Ecken geht) immer so, dass die Hypotenuse der Durchmesser des Kreises ist.

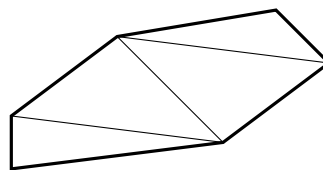
Der Satz wird durch die folgende Skizze illustriert. An der oberen Ecke des roten Dreiecks befindet sich der rechte Winkel.



★**Aufgabe 17.18.** Man kann sich leicht überlegen, warum der Satz des Thales wahr ist. Der gestrichelt eingezeichnete Radius in der obigen Skizze gibt einen Hinweis darauf, wie man vorgehen sollte. Vielleicht überlegen Sie erst einmal selbst, bevor Sie sich die Lösung anschauen.

Polygon

Wir wissen nun schon ein paar Dinge über rechtwinklige Dreiecke, aber was ist mit anderen geometrischen Figuren? Allgemeine Dreiecke kann man, wie wir bereits gesehen haben, in zwei rechtwinklige zerlegen. Aber auch beliebige *Polygone*¹⁸ kann man sich als zusammengesetzt aus Dreiecken vorstellen. Hier z. B. ein Sechseck, das in vier Dreiecke zerschnitten wurde:



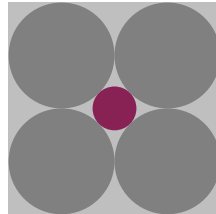
Grundsätzlich können Sie von jedem Polygon solange Dreiecke abschneiden, bis nur noch eines übrig bleibt.

Aufgabe 17.19. Zerlegen Sie ein Fünfeck in Dreiecke und schließen Sie daraus, welchen Wert die Summe der Innenwinkel eines Fünfecks hat.

Aufgabe 17.20. Verallgemeinern Sie die Aussage aus der letzten Aufgabe auf Polygone mit n Ecken.

¹⁸**Polygone** bzw. *Vielecke* sind geometrische Figuren, die durch gerade Linien begrenzt werden.

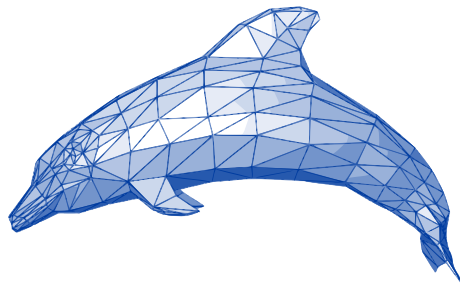
Aufgabe 17.21. In der Skizze unten sehen wir ein Quadrat, in dem vier gleich große (graue) Kreise liegen. Jeder dieser Kreise berührt zwei der anderen Kreise und zwei Seiten des Quadrats in jeweils genau einem Punkt. Der kleine rote Kreis berührt jeden der vier grauen Kreise in genau einem Punkt. Geben Sie den Radius des roten Kreises in Abhängigkeit von der Seitenlänge des Quadrats an.



Aufgabe 17.22. Und hier noch drei Aufgaben, die der letzten ähneln. Wenn sich Figuren berühren, dann ist es immer so gemeint, dass sie sich wie oben in genau einem Punkt berühren. Es geht jeweils darum, den Radius und die Lage des Mittelpunkts des kleinen roten Kreises in Abhängigkeit von den größeren grauen Quadraten bzw. Kreisen auszudrücken. (Stellen Sie sich vor, Sie wollten die Skizzen unten in einem Grafikprogramm nachstellen. Welche Werte müssten Sie diesem Programm übergeben?)



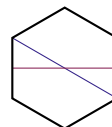
Die Möglichkeit der Zerlegung in Dreiecke wird heutzutage besonders intensiv von Grafikprozessoren genutzt. Moderne **GPUs** sind daraufhin optimiert, möglichst schnell möglichst viele Dreiecke zeichnen zu können. Polygone mit mehr als drei Ecken werden in Dreiecke zerlegt und andere geometrische Figuren werden durch Polygone approximiert. Hier ein simples Beispiel für ein solches *Mesh*:



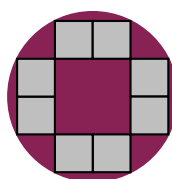
Sollten Sie mal vorhaben, eigene **Shader** zu programmieren, dann werden Sie mit Sicherheit alle geometrischen Kenntnisse aus diesem Abschnitt und den folgenden benötigen.

Wenn Sie aus den vorherigen Seiten nur eine wesentliche Aussage mitnehmen sollten, dann hoffentlich die, dass man sich bei fast allen geometrischen Problemen damit behelfen kann, (rechtwinklige) Dreiecke zu erkennen und deren Eigenschaften auszunutzen. Viele der folgenden Aufgaben liefern Beispiele dafür.

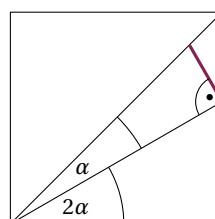
Aufgabe 17.23. In der Skizze sehen Sie ein regelmäßiges Sechseck bzw. *Hexagon*. Die blaue Verbindungslinie zweier Ecken hat die Länge 4. Wie lang ist die rote Linie, die rechtwinklig zu den Sechseckseiten verläuft, die sie berührt?



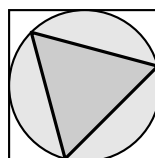
Aufgabe 17.24. In der folgenden Skizze haben alle acht Quadrate dieselbe Fläche und jedes Quadrat berührt den Kreis in genau einem Punkt. Quadrate, die sich berühren, haben entweder nur einen Eckpunkt oder eine ganze Seite gemein. Der Kreis hat den Radius 1. Was ist die Gesamtfläche der Quadrate?



Aufgabe 17.25. Die Länge der Diagonale in dem Quadrat unten ist 35. Was müsste man in einen Taschenrechner eingeben, um die Länge der rot eingezeichneten Strecke zu erhalten?

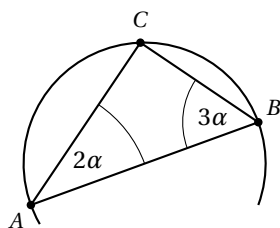


Aufgabe 17.26. Das Quadrat in der folgenden Skizze hat die Fläche 75 und berührt den Kreis genau mit seinen vier Seitenmittelpunkten. Das Dreieck ist gleichseitig und berührt den Kreis genau mit seinen drei Eckpunkten. Geben Sie die Seitenlänge des Dreiecks an. (Hinweis: Die Lösung ist rational, man braucht keine elektronische Hilfe.)

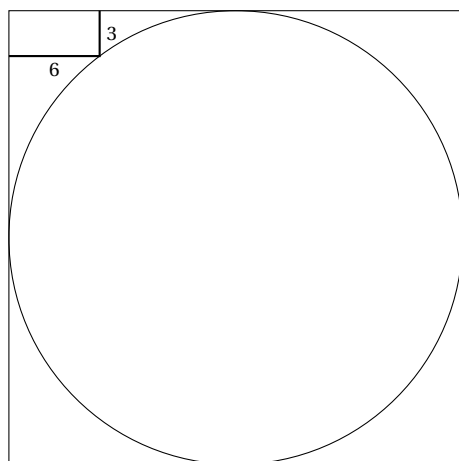


Aufgabe 17.27. Ein gleichschenkliges Dreieck hat die Seitenlängen 10, 13 und 13. Geben Sie seine Fläche an.

Aufgabe 17.28. Der Radius des Kreises unten ist 50. Die Strecke von A nach B geht durch den Mittelpunkt. Was müsste man in einen Taschenrechner eingeben, um die Länge der Strecke von B nach C zu ermitteln.



★**Aufgabe 17.29.** Welchen Radius hat der Kreis? (Wenn man die richtige Idee hat, ist es recht einfach. Man braucht eigentlich nur den Satz des Pythagoras.)



Bemerkungen

Es ist nach diesem Überblick sicher empfehlenswert, sich noch einmal die Ausführungen über Winkel im [Abschnitt über die komplexen Zahlen](#) anzuschauen.

In der Schule ist Ihnen vielleicht eingebläut worden, dass man immer so etwas wie „LE“ (für *Längeneinheiten*) oder „FE“ (für *Flächeneinheiten*) hinter die jeweiligen Zahlen schreiben solle. Das ist in der Mathematik jedoch gar nicht nötig. Es ist beispielsweise völlig OK zu sagen, eine Strecke habe die Länge 2. Die Angabe von Einheiten ist allerdings wichtig in den Natur- und Ingenieurwissenschaften.

Ein Einheitszeichen, das Sie jedoch auf keinen Fall vergessen sollten, ist das Gradzeichen! Wenn Winkelmaße ohne das Symbol $^\circ$ angegeben werden, dann ist grundsätzlich immer [Bogenmaß](#) gemeint.

Im Gegensatz zur Hypotenuse sind An- und Gegenkathete in einem rechtwinkligen Dreieck keine Bezeichnungen, die eine Seite eindeutig identifizieren. Man muss immer dazu sagen, welcher Winkel gemeint ist.

Für die Frage, wann man bei Sinus, Kosinus und Tangens Klammern verwendet, gibt es eine in der Praxis häufig verwendete, aber selten ausgesprochene Regel: Ohne Klammern erstreckt sich der „Wirkungsbereich“ bis zum nächsten Operator. In diesem Sinne ist $\sin \alpha \beta + c$ eine Abkürzung für $\sin(\alpha \beta) + c$, weil der Operator $+$ für

den Abschluss sorgt, während für das Produkt $\alpha\beta$ kein Rechenzeichen verwendet wurde. Wenn Sie selbst so etwas aufschreiben und sich nicht sicher sind, so sollten Sie im Zweifelsfall lieber ein Klammernpaar zu viel als eines zu wenig benutzen. Sie müssen aber damit leben, dass sie in der Literatur auf solche Schreibweisen stoßen werden.

18. Die trigonometrischen Funktionen

In diesem Abschnitt werden wir uns die wichtigen Funktionen Sinus, Kosinus und Tangens noch etwas genauer ansehen. Dafür benötigen wir zunächst das in Mathematik, Informatik und Naturwissenschaften verwendete Winkelmaß und deshalb müssen wir zuerst ein bisschen über Kreise nachdenken.

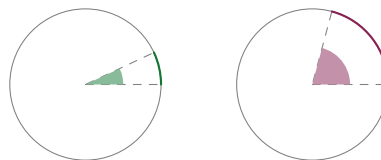
Es ist anschaulich klar und war daher schon den alten Hochkulturen bekannt, dass der **Umfang** eines Kreises unabhängig von seiner Größe immer dasselbe Verhältnis zu seinem **Durchmesser** hat. Dieses Verhältnis – die sogenannte **Kreiszahl** – nennt man π , allerdings erst seit dem 18. Jahrhundert. Der Name π („pi“) ist der griechische Buchstabe, der am ehesten dem lateinischen p entspricht. Er wurde gewählt, weil er der Anfangsbuchstabe von Wörtern wie **Peripherie** und **Perimeter** ist.

Satz 18.1

Ist U der Umfang eines Kreises und d sein Durchmesser, so gilt $U = \pi d$. Anders ausgedrückt gilt $U = 2\pi r$, wenn r der Radius des Kreises ist.

Seit der Antike weiß man, dass π *ungefähr* den Wert 3.14 hat. Dieser Näherungswert wurde im Laufe der Jahrhunderte immer weiter verbessert.¹⁹ Inzwischen (Stand 2025) hat man mithilfe von Computern 300 Billionen Nachkommastellen berechnet. Solche Rekorde sind aber eher für die Informatik interessant. Für Anwendungen reichen die 15 Nachkommastellen, die eine typische Programmiersprache zur Verfügung stellt, völlig aus. Und für die Mathematik ist entscheidend, dass π **irrational** ist.²⁰

Was ebenfalls anschaulich klar sein sollte: Die Länge eines **Kreisbogens** ist proportional zum überstrichenen Winkel.²¹ Mit anderen Worten ist die rote Linie in der folgenden Skizze genau dann dreimal so lang wie die grüne, wenn der rote Winkel dreimal so groß wie der grüne ist.



¹⁹Mehr dazu in dem Video [Eine kurze Geschichte der Kreiszahl Pi](#).

²⁰Siehe dazu Fußnote 13 auf Seite 96.

²¹Man kann das auch ohne Verweis auf die Anschauung streng mathematisch begründen, aber dafür braucht man die **Integralrechnung** und muss insbesondere erst einmal festlegen, was mit der **Länge** einer beliebigen Linie überhaupt gemeint ist. Für unsere Zwecke reicht es, sich die betreffende Länge als den Weg vorzustellen, den eine winzige Ameise auf dem Kreisbogen zurücklegt.

Aufgabe 18.1. Wie lang ist ein Kreisbogen, der einer halben Umdrehung um den Kreis entspricht, bei einem Kreis mit dem Radius 1?

Die Länge des entsprechenden Kreisbogens im Einheitskreis nimmt man nun als Maß für Winkel und nennt das *Radian* oder *Bogenmaß*. Im Gegensatz zum historisch bedingten Gradmaß ist diese Einteilung nicht willkürlich. Mathematisch ist sie aus verschiedenen Gründen die einzig sinnvolle.

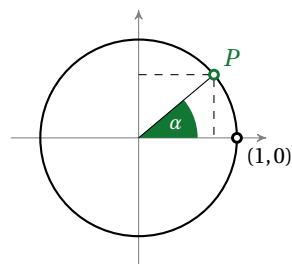
Radian / Bogenmaß

Zum Glück ist die Umrechnung zwischen beiden Maßen ganz einfach. Man muss sich lediglich einen der Werte aus der folgenden Tabelle merken. Der Rest ist simpler Dreisatz.

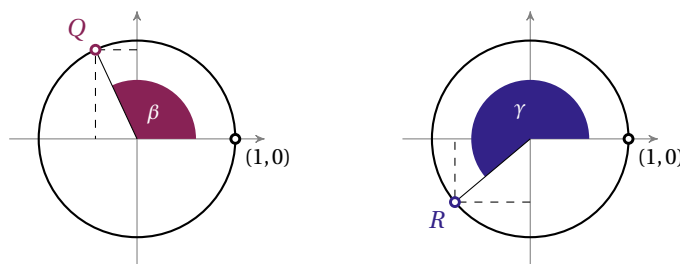
Bogenmaß	$\pi/6$	$\pi/4$	$\pi/2$	π	2π
Grad	30°	45°	90°	180°	360°

Beachten sollte man dabei allerdings, dass für Grad das Einheitenzeichen $^\circ$ verwendet wird, während das Bogenmaß *einfach nur eine Zahl* ist. Taschenrechner bieten in der Regel Tasten wie **RAD** und **DEG** für das Umschalten zwischen Radian und Grad (engl. *degree*) an. Die meisten Programmiersprachen erwarten, dass Winkel in Bogenmaß angegeben werden.

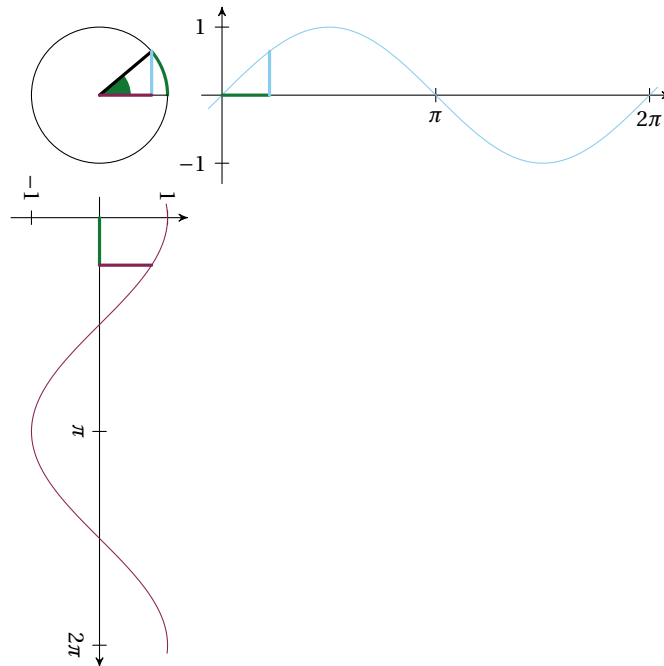
Durch die Einführung des Bogenmaßes ergibt sich im *kartesischen Koordinatensystem* im Einheitskreis die folgende Situation:



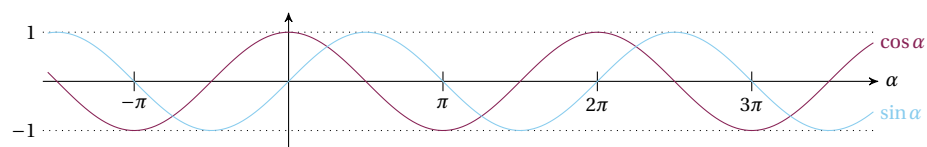
Der grüne Punkt P hat die Koordinaten $(\cos \alpha, \sin \alpha)$. Bisher ergibt das jedoch lediglich Sinn für Winkel zwischen 0 und $\pi/2$, weil wir Sinus und Kosinus über rechtwinklige Dreiecke definiert haben. Man kann das jedoch auf eine natürliche Art und Weise erweitern, indem man für *beliebige* Winkel α Kosinus und Sinus so definiert, dass $(\cos \alpha, \sin \alpha)$ die Koordinaten des Punktes auf dem Einheitskreis sind, der ausgehend vom Punkt $(1,0)$ um den Winkel α (in mathematisch positiver Drehrichtung) gedreht wurde.



Q und R in der obigen Skizze haben dann also die Koordinaten $(\cos \beta, \sin \beta)$ bzw. $(\sin \beta, \cos \beta)$. Betrachtet man Sinus und Kosinus nun als **Funktionen**, die Winkeln (in Bogenmaß) Zahlen zuordnen, so erhält man die aus der Schule bekannten Funktionsgraphen, wenn man den Weg eines Punktes um den Einheitskreis herum verfolgt und dabei dessen Winkel zusammen mit seiner x - bzw. y -Koordinate protokolliert. Das zeigt die folgende Skizze.



Wenn wir wie im letzten Abschnitt auch Winkel zulassen, die negativ oder größer als 2π (also 360°) sind, erhalten wir zwei Funktionen, die auf ganz $\mathbb{R}$ definiert sind.



Wir können an diesen Kurven einige wesentliche Eigenschaften von Sinus und Kosinus sowie wichtige Zusammenhänge zwischen beiden Funktionen erkennen:

- Durch die entsprechende Behandlung von Winkeln außerhalb des Intervalls $[0, 2\pi)$ ergibt sich, dass sich nach einer kompletten Umdrehung alles wiederholt. In Formeln ausgedrückt bedeutet das, dass

$$\sin(x + 2k\pi) = \sin x \quad \text{sowie} \quad \cos(x + 2k\pi) = \cos x$$

für alle $k \in \mathbb{Z}$ und alle $x \in \mathbb{R}$ gelten. Man sagt auch, dass beide Funktionen 2π -periodisch sind.

- Der Kosinus hat beim Weg um den Kreis herum quasi einen „Vorsprung“ von einem rechten Winkel gegenüber dem Sinus. Jeder Wert, den der Kosinus

periodische Funktion

annimmt, wird „90 Grad später“ vom Sinus angenommen.²² Mit anderen Worten gilt $\cos x = \sin(x + \pi/2)$ und umgekehrt $\sin x = \cos(x - \pi/2)$ für alle $x \in \mathbb{R}$.

- Wenn man den Kosinus bei der Kreisumdrehung betrachtet, fällt einem auf, dass man dieselben Werte erhielte, wenn man im Uhrzeigersinn drehen würde. Deshalb ist der Graph der Kosinusfunktion symmetrisch zur y -Achse; es gilt also $\cos(-x) = \cos x$ für alle $x \in \mathbb{R}$. Funktionen mit dieser Eigenschaft nennt man *gerade*. Der Sinus ist hingegen punktsymmetrisch zum Ursprung des Koordinatensystems, d. h., es gilt $\sin(-x) = -\sin x$ für alle x . Solche Funktionen nennt man *ungerade*.
- Wenn man die Hypotenuse um die Hälfte des Kreisumfangs weiterdreht, drehen sich offenbar die Richtungen der Katheten um. Das lässt sich so ausdrücken (wiederum für alle $x \in \mathbb{R}$):

$$\begin{aligned}\cos(x + \pi) &= \cos(x - \pi) = -\cos x \\ \sin(x + \pi) &= \sin(x - \pi) = -\sin x\end{aligned}$$

- Der Punkt mit dem Winkel $x + \pi/2$ liegt auf dem Kreis direkt gegenüber dem Punkt mit dem Winkel $x - \pi/2$. Daraus ergibt sich:

$$\begin{aligned}\cos(x + \pi/2) &= -\cos(x - \pi/2) \\ \sin(x + \pi/2) &= -\sin(x - \pi/2)\end{aligned}$$

- Sinus und Kosinus können nun auch negative Werte annehmen. Der Sinus ist genau dann positiv, wenn sich die Hypotenuse bei der Umdrehung in der oberen Hälfte des Kreises befindet. Der Kosinus ist genau dann positiv, wenn sie sich in der rechten Hälfte des Kreises befindet.²³
- Da die Hypotenuse die längste Seite im Dreieck ist, können Sinus und Kosinus nur Werte annehmen, deren Betrag nicht größer als 1 ist. Die Extremwerte 0, 1 und -1 werden immer dann angenommen, wenn man es mit einem „entarteten“ Dreieck zu tun hat, wenn also die Hypotenuse genau auf einer der Achsen liegt.
- Aus dem **Satz des Pythagoras** folgt, dass $\sin^2 x + \cos^2 x = 1$ für alle $x \in \mathbb{R}$ gilt. Dabei verwenden wir die allgemein übliche Konvention, dass $\sin^k x$ eine Abkürzung für $(\sin x)^k$ ist.

gerade Funktion

ungerade Funktion

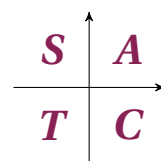
 $\sin^k x$

Aufgabe 18.2. Aus den obigen Bemerkungen folgt insbesondere, wo die **Nullstellen** von Sinus und Kosinus sind. Nämlich wo?

Wir fassen diese Überlegungen der Übersichtlichkeit halber in dem folgenden Lemma zusammen.

²²In der Physik sagt man, dass die beiden Funktionen *phasenverschoben* sind.

²³Dafür gibt es die Eselsbrücke **All Science Teachers are Crazy**. Trägt man diese Buchstaben wie unten gegen den Uhrzeigersinn auf, so erhält man, dass im ersten Quadranten alle drei Funktionen positiv sind, in den folgenden nur **s**in, nur **c**os bzw. nur **t**an.



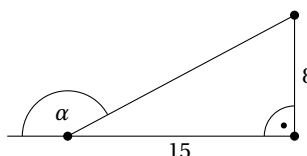
Lemma 18.2**Grundlegende Eigenschaften von Sinus und Kosinus**

Sinus und Kosinus sind zwei Funktionen mit Definitionsbereich $\mathbb{R}$ und Wertebereich $[-1, 1]$. Für alle $x \in \mathbb{R}$ und alle $k \in \mathbb{Z}$ gilt:

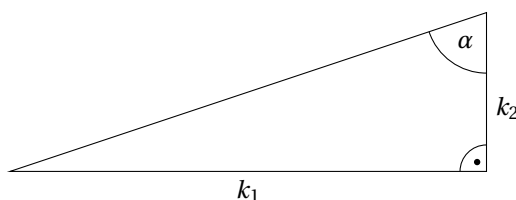
$$\begin{aligned}\sin(x + 2k\pi) &= \sin x \\ \cos(x + 2k\pi) &= \cos x \\ \sin(x + \frac{\pi}{2}) &= \cos x \\ \cos(x - \frac{\pi}{2}) &= \sin x \\ \cos(-x) &= \cos x \\ \sin(-x) &= -\sin x \\ \cos(x + \pi) &= \cos(x - \pi) = -\cos x \\ \sin(x + \pi) &= \sin(x - \pi) = -\sin x \\ \cos(x + \frac{\pi}{2}) &= -\cos(x - \frac{\pi}{2}) \\ \sin(x + \frac{\pi}{2}) &= \sin(x - \frac{\pi}{2}) \\ \cos(\frac{\pi}{2} + k\pi) &= \sin(k\pi) = 0 \\ \cos(k\pi) &= \sin(\frac{\pi}{2} + k\pi) = (-1)^k \\ \sin^2 x + \cos^2 x &= 1\end{aligned}$$

Außer den angegebenen Werten gibt es keine weiteren Stellen, an denen Sinus oder Kosinus verschwinden oder Extremwerte annehmen.

Aufgabe 18.3. Geben Sie für die folgende Skizze den Sinus von α an. (Hinweis: Das Ergebnis ist ein Bruch.)



Aufgabe 18.4. Sie haben als Werbegeschenk einen billigen Taschenrechner erhalten, der nur die vier Grundrechenarten beherrscht. Wie könnten Sie damit für das folgende Dreieck $\sin^2 \alpha$ berechnen, wenn Sie k_1 und k_2 kennen?

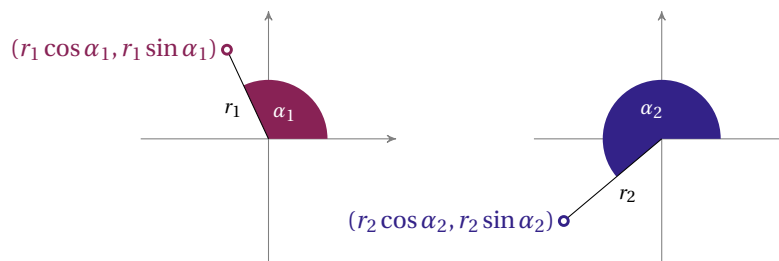


Aufgabe 18.5. Bei dieser Aufgabe beherrscht Ihr Taschenrechner die trigonometrischen Funktionen und hat auch eine Taste für π , allerdings ist er defekt. Erstens kann man ihn nicht mehr auf Gradmaß umstellen und muss daher immer in Bogenmaß

rechnen. Zweitens funktioniert die Taste für die Sinusfunktion gar nicht mehr. Und drittens arbeitet zu allem Überfluss die Kosinusfunktion nur für Winkel im Intervall $[0, \pi/2]$ korrekt. Finden Sie mithilfe von Lemma 18.2 heraus, was man eingeben muss, um trotzdem die folgenden Werte korrekt zu berechnen:

$$\begin{array}{cccc} \cos(2\pi/3) & \cos(5\pi/4) & \cos(-\pi/3) & \cos(8\pi + 1) \\ \sin(2\pi/5) & \sin(7\pi/4) & \cos(20^\circ) & \sin(250^\circ) \end{array}$$

Eine ganz wesentliche Einsicht ist, dass der Zusammenhang zwischen Winkeln und kartesischen Koordinaten, den wir auf Seite 155 hergestellt haben, nicht nur für Punkte auf dem Einheitskreis gilt, sondern für *alle* Punkte der Ebene. Das liegt daran, dass sich durch maßstabsgetreue Vergrößerung und Verkleinerung Winkel nicht ändern. Anschaulich kann man sich das folgendermaßen klarmachen: Wenn man vom Ursprung des Koordinatensystems zu einem bestimmten Punkt kommen will, dann kann man den Weg dorthin auf zwei Arten beschreiben. Entweder gibt man an, dass man sich x Einheiten nach rechts bewegt und anschließend y Einheiten nach oben. (Dabei bedeuten negative Zahlen, dass man nach links bzw. unten wandert.) Oder man dreht sich zuerst um den Winkel α (wobei man vor der Drehung nach rechts schaute) und wandert dann in dieser Richtung r Einheiten, wobei r der Abstand des Punktes vom Ursprung ist. In beiden Fällen hat man durch die Angabe von zwei Zahlen – x und y bzw. r und α – die Lage des Punktes eindeutig beschrieben. Der Zusammenhang zwischen diesen beiden Angaben ist $(x, y) = (r \cos \alpha, r \sin \alpha)$. Für $r = 1$ erhält man die Punkte auf dem Einheitskreis, aber man kann für beliebige positive²⁴ Zahlen r so verfahren. Zwei Beispiele zur Illustration:



Gibt man die Lage eines Punktes auf diese Art mittels Abstand und Winkel an, so nennt man diese beiden Werte die *Polarkoordinaten* des Punktes. Das führt auch zu einer entsprechenden Darstellung der *komplexen Zahlen*. Die Begriffe hatten wir bereits eingeführt, mit den trigonometrischen Funktionen können wir nun den folgenden Zusammenhang herstellen: Hat eine komplexe Zahl z den *Betrag* r und das *Argument* α so sind ihr *Real- und Imaginärteil* $r \cos \alpha$ und $r \sin \alpha$:

$$z = r(\cos \alpha + i \sin \alpha)$$

Da wir uns bereits überlegt hatten, wie die komplexe Multiplikation geometrisch interpretiert werden kann, können wir über diesen Umweg eine weitere wichtige

Polarkoordinaten

²⁴Für $r = 0$ – also für den Ursprung $(0, 0)$ – ergibt die Angabe eines Winkels offenbar keinen Sinn.

Eigenschaft der trigonometrischen Funktionen herleiten. Sind z_1 und z_2 zwei komplexe Zahlen mit $z_i = r_i(\cos \alpha_i + i \sin \alpha_i)$, dann erhält man einerseits algebraisch durch einfaches Ausmultiplizieren

$$\begin{aligned} z_1 z_2 &= (r_1 \cos \alpha_1 + i r_1 \sin \alpha_1) \cdot (r_2 \cos \alpha_2 + i r_2 \sin \alpha_2) \\ &= r_1 r_2 ((\cos \alpha_1 \cos \alpha_2 - \sin \alpha_1 \sin \alpha_2) + i(\sin \alpha_1 \cos \alpha_2 + \cos \alpha_1 \sin \alpha_2)) \end{aligned}$$

und andererseits geometrisch (siehe Seite 112)

$$z_1 z_2 = r_1 r_2 (\cos(\alpha_1 + \alpha_2) + i \sin(\alpha_1 + \alpha_2)).$$

Vergleicht man die beiden Werte, so ergibt sich der folgende Satz:

Satz 18.3

Additionstheoreme für Sinus und Kosinus

Für alle $x, y \in \mathbb{R}$ gilt:

$$\begin{aligned} \cos(x \pm y) &= \cos x \cos y \mp \sin x \sin y \\ \sin(x \pm y) &= \sin x \cos y \pm \sin y \cos x \end{aligned}$$

Dabei erhält man die beiden Formeln für $x - y$, indem man in die für $x + y$ jeweils $-y$ statt y einsetzt und dann Lemma 18.2 anwendet.

Im vorherigen Abschnitt hatten wir bereits geometrisch die Werte der trigonometrischen Funktionen für ein paar wichtige Winkel hergeleitet. Die werden in der folgenden Tabelle noch einmal zusammengefasst. Dabei sind die beiden grünen Zeilen am Ende lediglich als Merkhilfe gedacht.

α	0	$\pi/6$	$\pi/4$	$\pi/3$	$\pi/2$
	0°	30°	45°	60°	90°
$\sin \alpha$	0	1/2	$1/\sqrt{2}$	$\sqrt{3}/2$	1
$\cos \alpha$	1	$\sqrt{3}/2$	$1/\sqrt{2}$	1/2	0
n	0	1	2	3	4
$\sqrt{n}/2$	0	1/2	$1/\sqrt{2}$	$\sqrt{3}/2$	1

Aufgabe 18.6. Berechnen Sie mit dem ersten der beiden Additionstheoreme $\cos 2x$ und ersetzen Sie im Ergebnis $\sin^2 x$ mithilfe der letzten Formel von Lemma 18.2 durch einen Ausdruck, in dem der Sinus nicht mehr vorkommt. Lösen Sie das Ergebnis nach $\cos x$ auf.

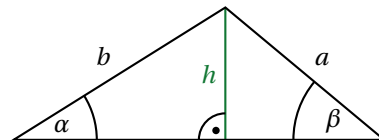
Aufgabe 18.7. Verwenden Sie die Lösung von Aufgabe 18.6, um $\cos(\pi/8)$ ohne elektronische Hilfen zu bestimmen.

Es folgen zwei Aussagen über beliebige Dreiecke, bei denen Sinus und Kosinus eine wichtige Rolle spielen.

Sinussatz

In jedem Dreieck entspricht das Verhältnis zweier Seiten dem Verhältnis der Sinuswerte der gegenüberliegenden Seiten. Als Formel ausgedrückt bedeutet das, dass in der folgenden Skizze diese Beziehung gilt:

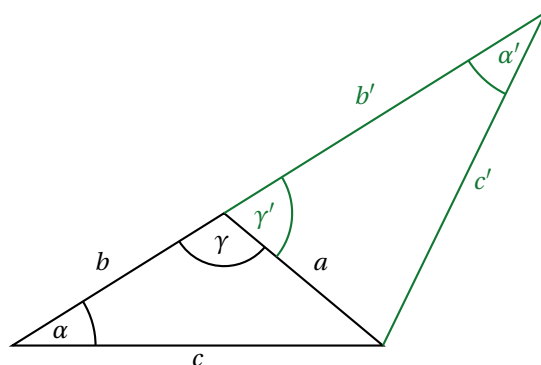
$$\frac{a}{b} = \frac{\sin \alpha}{\sin \beta}$$



Der Sinussatz liefert gewissermaßen die nachträgliche Begründung für die Aussage, dass in rechtwinkligen Dreiecken die Hypotenuse als Seite gegenüber dem rechten Winkel immer die längste Seite ist: Weil der Sinus des rechten Winkels den maximalen Wert 1 hat, kann keine andere Seite länger sein.

Die Skizze zeigt – zumindest für den Fall, dass α und β spitze Winkel sind – durch die Zerlegung in zwei rechtwinklige Dreiecke auch gleich, warum die Aussage wahr ist; Im linken rechtwinkligen Dreieck gilt $\sin \alpha = h/b$, im rechten $\sin \beta = h/a$. Dividiert man $\sin \alpha$ durch $\sin \beta$, so kürzt sich h heraus und man erhält sofort die gesuchte Beziehung.

Dass der Sinussatz ebenso für stumpfe Winkel gilt, zeigt die folgende Skizze, in der die Seite b verlängert wurde und c' so lang wie c sein soll.



Da γ stumpf ist, muss γ' spitz sein. Im grünen Dreieck gilt also der Sinussatz und damit haben wir:

$$\frac{\sin \alpha'}{\sin \gamma'} = \frac{a}{c'} = \frac{a}{c} \quad (18.1)$$

Da das gesamte, sich von α nach α' erstreckende Dreieck gleichschenkelig ist, muss nach Aufgabe 17.5 $\alpha = \alpha'$ gelten. Außerdem gilt $\gamma + \gamma' = \pi$, weil γ und γ' Ergänzungswinkel sind, und daher $\sin \gamma' = \sin(\pi - \gamma) = \sin \gamma$ nach Lemma 18.2. Setzt man das beides in (18.1) ein, so erhält man die gewünschte Beziehung.

Satz 18.4

Satz 18.5**Kosinussatz**

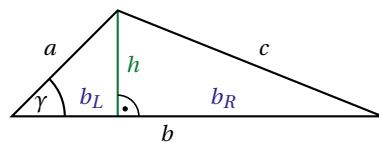
In jedem Dreieck gilt

$$c^2 = a^2 + b^2 - 2ab \cos \gamma,$$

wenn a , b und c die Seiten des Dreiecks sind und γ der Winkel gegenüber der Seite c ist.

Der Kosinussatz ist eine Verallgemeinerung des [Satzes des Pythagoras](#), denn für den Spezialfall, dass γ ein rechter Winkel ist, verschwindet der dritte Term rechts vom Gleichheitszeichen. Ansonsten handelt es sich um eine Art „Korrekturterm“, der c gewissermaßen größer oder kleiner macht, je nachdem, ob γ größer oder kleiner als $\pi/2$ ist.

Zur Begründung zerlegen wir das Dreieck folgendermaßen in zwei rechtwinklige:



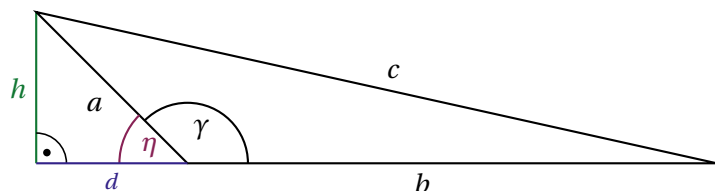
Für die Länge von c gilt nach Pythagoras im rechten Dreieck:

$$c^2 = h^2 + b_R^2 \quad (18.2)$$

Die Seite b_R ergibt sich als Differenz $b - b_L$. Im linken Dreieck erhält man die Gleichung $b_L = a \cdot \cos \gamma$ und für h gilt nach Pythagoras $h^2 = a^2 - b_L^2$. Setzt man in Gleichung (18.2) ein, so ergibt sich:

$$\begin{aligned} c^2 &= h^2 + b_R^2 = a^2 - b_L^2 + (b - b_L)^2 \\ &= a^2 - b_L^2 + b^2 - 2bb_L + b_L^2 = a^2 + b^2 - 2bb_L \\ &= a^2 + b^2 - 2ab \cos \gamma \end{aligned}$$

Die obige Konstruktion funktioniert jedoch nur für spitze Winkel. Ist γ stumpf, so fügen wir ein rechtwinkliges Dreieck hinzu:



Im großen rechtwinkligen Dreieck mit der Hypotenuse c gilt $c^2 = h^2 + (b + d)^2$. Im hinzugefügten rechtwinkligen Dreieck lesen wir $d = a \cos \eta$ ab und mit Pythagoras außerdem $h^2 = a^2 - d^2$. Fügen wir alles zusammen, so erhalten wir:

$$c^2 = h^2 + (b + d)^2 = a^2 - d^2 + b^2 + 2bd + d^2$$

$$= a^2 + b^2 + 2bd = a^2 + b^2 + 2ab \cos \eta \quad (18.3)$$

Wegen $\eta + \gamma = \pi$ folgt aber mit Lemma 18.2 $\cos \eta = \cos(\pi - \gamma) = -\cos \gamma$, so dass sich aus dem letzten Plus in (18.3) ein Minus wird.

Aus dem Kosinussatz folgt nebenbei auch die folgende Aussage:

Umgekehrter Satz des Pythagoras

Wenn a , b und c die Seiten eines Dreiecks sind und $c^2 = a^2 + b^2$ gilt, dann ist das Dreieck rechtwinklig mit Hypotenuse c .

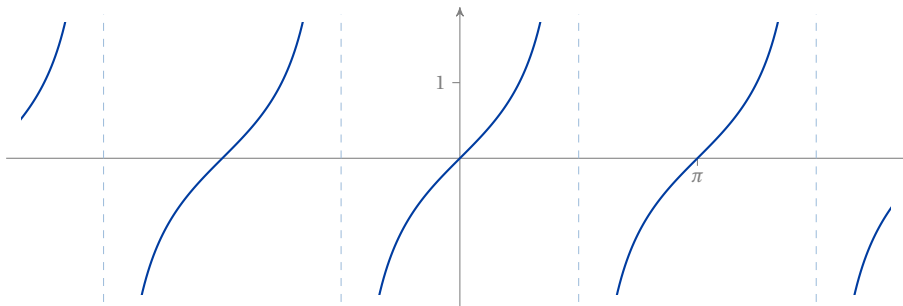
Korollar 18.6

Die Begründung ist, dass nach dem Kosinussatz $c^2 = a^2 + b^2 - 2ab \cos \gamma$ für den Winkel γ gegenüber der Seite c gilt. Nach Voraussetzung muss $2ab \cos \gamma$ und damit $\cos \gamma$ jedoch verschwinden. Das geht nur, wenn γ ein rechter Winkel ist.

Aufgabe 18.8. Ist es möglich, positive Zahlen a , b und c anzugeben, die nicht die Seiten eines Dreiecks sein können? Ist das auch möglich, wenn $a^2 + b^2 = c^2$ gilt?

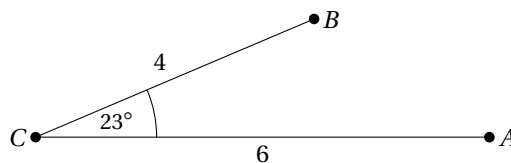
Bisher fehlte noch der Tangens. Wie im letzten Abschnitt definieren wir den nun ganz allgemein durch $\tan x = \sin x / \cos x$. Allerdings müssen wir hier etwas aufpassen, weil das nicht für alle reellen Zahlen x möglich ist, sondern nur dann, wenn der Kosinus nicht verschwindet.

Da Sinus und Kosinus nach Lemma 18.2 beide periodisch, aber gegeneinander phasenverschoben sind, ist der Tangens periodisch mit der „halben“ Periode π . Der Funktionsgraph sieht so aus wie in der folgenden Grafik.



Denkt man in Polarkoordinaten, so ist der Tangens von α also der Quotient von y - und x -Koordinate des Punktes $(r \cos \alpha, r \sin \alpha)$. Und für Punkte auf der y -Achse ist der Tangens nicht definiert.

Aufgabe 18.9. Berechnen Sie mit dem Kosinussatz und einem Taschenrechner oder Smartphone den Abstand der Punkte A und B in der folgenden Skizze.



Bemerkungen

Das Arbeiten mit dem falschen Winkelmaß ist eine typische Fehlerquelle! Oft wird beim Einsatz von Taschenrechnern oder beim Schreiben von Programmen vergessen, dass Winkel in Bogenmaß angegeben werden sollten.

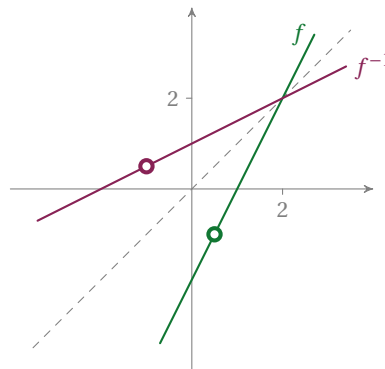
Es bringt nichts, die Formeln aus Lemma 18.2 alle auswendig zu lernen. Dafür sehen sie zu ähnlich aus und man kommt nur durcheinander. Sie sollten stattdessen die Skizze auf Seite 156 im Kopf haben und den ungefähren Verlauf der Funktionsgraphen von Sinus und Kosinus. Damit kann man sich in der Regel Umformungen „herleiten“, wenn man sie mal braucht.

19. Die Arkusfunktionen

Wir wissen bereits, wie man aus dem Abstand eines Punktes vom Ursprung und dem zugehörigen Winkel – also aus seinen **Polarkoordinaten** – die kartesischen Koordinaten berechnet. Aber wie kommt man umgekehrt von den kartesischen zu den Polarkoordinaten? Dafür müssen wir uns noch etwas ausführlicher mit Umkehrfunktionen beschäftigen.

Aufgabe 19.1. Ein Teil der eben aufgeworfenen Frage lässt sich jedoch ganz leicht beantworten. Welchen Abstand hat der Punkt (x, y) vom Ursprung?

Für Funktionen, die reelle Zahlen auf reelle Zahlen abbilden, kann man nicht nur einen **Funktionsgraphen** zeichnen, man kann auch direkt den Funktionsgraphen der **Umkehrfunktion** „sehen“. Schauen wir uns dafür ein Beispiel an. Die Funktion $f(x) = 2x - 2$ hat die Umkehrfunktion $f^{-1}(y) = y/2 + 1$. Beide sieht man in der folgenden Grafik.

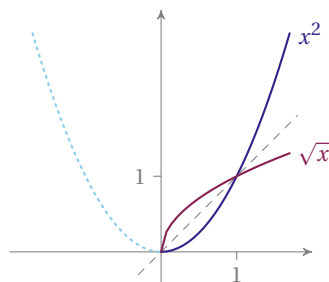


Hervorgehoben sind der Punkt $(1/2, -1)$ auf dem Funktionsgraphen von f und der Punkt $(-1, 1/2)$ auf dem Funktionsgraphen von f^{-1} . Die Zahlenwerte sind dabei irrelevant. Entscheidend ist, dass es zu jedem Punkt (x, y) des einen Funktionsgraphen den entsprechenden Punkt (y, x) auf dem anderen geben muss. Das war ja gerade die Definition der Umkehrfunktion. Geometrisch entsteht der Graph der Umkehrfunktion einer Funktion f also einfach dadurch, dass man den Funktionsgraphen von f an der (gestrichelt eingezeichneten) **Hauptdiagonalen** – dem Funktionsgraphen der **Identität** $\text{id}_{\mathbb{R}}$ – spiegelt.

Hauptdiagonale

Aufgabe 19.2. Schauen Sie sich noch einmal Aufgabe 16.7 und Aufgabe 16.11 an und machen Sie sich klar, dass durch die hier beschriebene Spiegelung horizontale Geraden zu vertikalen werden und umgekehrt. Das ist gewissermaßen eine Illustration von Satz 16.1.

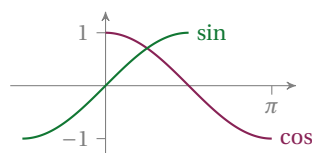
Hat man nun eine Funktion, die *nicht* injektiv ist, so kann man diese eventuell dadurch „injektiv machen“, dass man an ihrer statt nur einen *Teil* der Funktion betrachtet: Man verkleinert den Definitionsbereich, so dass mehrfach auftretende Funktionswerte verschwinden. Die folgende Grafik zeigt das am Beispiel der nicht injektiven Funktion $x \mapsto x^2$.



Lässt man die linke Hälfte der Parabel weg, so erhält man eine injektive Funktion, die immer noch x auf x^2 abbildet, deren Definitionsbereich nun aber $[0, \infty)$ statt $\mathbb{R}$ ist. Diese Funktion kann man umkehren und ihre Umkehrfunktion ist die bekannte Wurzelfunktion. (Dass man es so macht, ist übrigens lediglich die übliche Konvention. Man hätte auch die rechte Hälfte „ausblenden“ können. Das hätte ebenfalls zu einer injektiven Funktion geführt und die Umkehrfunktion hätte dann x auf $-\sqrt{x}$ abgebildet.)

Zurück zur Geometrie: Die trigonometrischen Funktionen Sinus, Kosinus und Tangens bekommen als Argumente – also als Eingaben – Winkel und geben als Funktionswerte Seitenverhältnisse aus. Wenn wir aus kartesischen Koordinaten Polarkoordinaten machen wollen, haben wir Seitenverhältnisse und wollen Winkel haben. Wir müssen die trigonometrischen Funktionen daher umkehren. Das geht jedoch so ohne Weiteres nicht. An den Funktionsgraphen auf den vorherigen Seiten sieht man, dass weder Sinus noch Kosinus noch Tangens injektiv sind, weil alle drei periodisch sind.

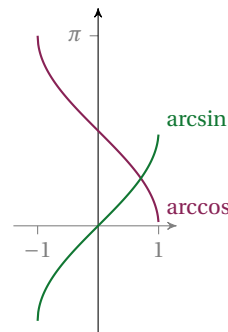
Um eine Umkehrfunktion zu erhalten, schränkt man daher so wie beim Beispiel der Parabel die Definitionsbereiche ein. Für den Kosinus nimmt man $[0, \pi]$ und für den Sinus $[-\pi/2, \pi/2]$.²⁵



²⁵Man hätte in beiden Fällen natürlich auch andere Stücke der Funktion nehmen können. Dass man gerade diese nimmt, ist lediglich eine Konvention.

Arkussinus /
Arkuskosinus

Durch Spiegeln an der Hauptdiagonalen erhält man nun die Umkehrfunktionen,²⁶ die man *Arkussinus* und *Arkuskosinus* nennt und für die man $\arcsin$ bzw. $\arccos$ schreibt. Der Präfix *Arkus* kommt vom lateinischen Wort für *Bogen*, weil diese Funktionen vom Seitenverhältnis zurück ins Bogenmaß rechnen.



Wie man an der Grafik sieht, ist der Definitionsbereich beider Funktion das Intervall $[-1, 1]$. Werte die kleiner als -1 oder größer als 1 sind, kann man also nicht einsetzen. Das ist nicht verwunderlich, weil solche Zahlen nicht als Funktionswerte von Sinus oder Kosinus vorkommen können.

Auf Taschenrechnern sind übrigens auch Bezeichnungen wie $\boxed{\sin^{-1}}$ statt $\arcsin$ üblich. Das kann allerdings zu Verwechslungen mit der bereits eingeführten Schreibweise $\sin^k x$ führen. Es ist nicht ohne Weiteres klar, ob mit $\sin^{-1} \alpha$ der Wert $\arcsin \alpha$ oder stattdessen $(\sin \alpha)^{-1}$ gemeint ist.

Da wir nun aus Seitenverhältnissen wieder Winkel machen können, können wir uns beispielsweise einen Punkt (x, y) vornehmen, seinen Abstand r vom Ursprung berechnen (siehe Aufgabe 19.1) und aus der Gleichung $\cos \alpha = x/r$ mit dem Arkuskosinus $\alpha = \arccos(x/r)$ machen. Aber leider liefert der Arkuskosinus nur Werte zwischen 0 und π . Das bedeutet, dass diese „Lösung“ für Punkte unterhalb der x -Achse falsche Angaben liefert. Man muss eine Fallunterscheidung machen:²⁷

$$\alpha = \begin{cases} \arccos(x/r) & y \geq 0 \\ -\arccos(x/r) & y < 0 \end{cases}$$

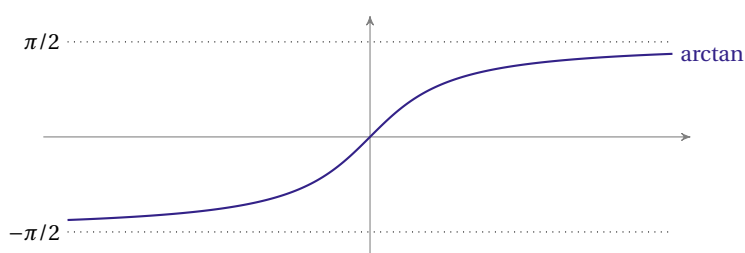
Aufgabe 19.3. Man kann das besagte Problem auch nicht aus der Welt schaffen, wenn man stattdessen den Zusammenhang $\sin \alpha = y/r$ und den Arkussinus verwendet. Wie sieht dann die Fallunterscheidung aus?

Arkustangens

Den Tangens kann man ebenfalls umkehren, nachdem man ihn vorher eingeschränkt hat. Die übliche Konvention ist, den Bereich zwischen $-\pi/2$ und $\pi/2$ auszuwählen. Wie Sie sich schon gedacht haben, heißt diese Umkehrfunktion *Arkustangens* und wird $\arctan$ geschrieben. Der Definitionsbereich dieser Funktion ist ganz $\mathbb{R}$.

²⁶Das ist natürlich leicht gesagt. Es ist zunächst überhaupt nicht klar, wie man diese Werte *berechnen* soll. Darum wollen wir uns aber im Moment nicht kümmern.

²⁷Für $r = 0$ ist der Ausdruck nicht definiert, aber das brauchen wir auch nicht.

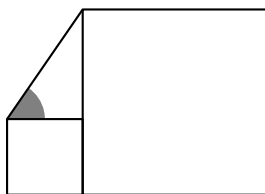


Wie die anderen beiden Arkusfunktionen kann auch der Arkustangens nicht den gesamten Kreis abdecken. Beim Umrechnen von Polarkoordinaten in kartesische bietet er jedoch den Vorteil, dass man nicht erst den Betrag berechnen muss. (Dafür muss man allerdings die Definitionslücken beachten.) Den Winkel α zum Punkt (x, y) erhält man durch

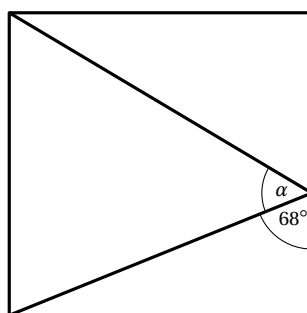
$$\alpha = \begin{cases} \arctan(y/x) & x > 0 \\ \arctan(y/x) + \pi & x < 0 \wedge y \neq 0 \\ \pi/2 & x = 0 \wedge y > 0 \\ -\pi/2 & x = 0 \wedge y < 0 \end{cases},$$

wenn es sich nicht um den Ursprung handelt. Siehe dazu auch die Bemerkung zur Hilfsfunktion `atan2` im folgenden Abschnitt über Trigonometrie im Computer.

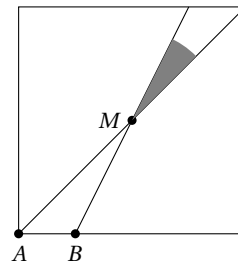
Aufgabe 19.4. Die Fläche des rechten Quadrats ist genau sechsmal so groß wie die Fläche des linken Quadrats. Wie berechnet man den grau eingezeichneten Winkel?



Aufgabe 19.5. Das äußere Viereck ist ein Quadrat. Wie kann man α ermitteln?



Aufgabe 19.6. In der folgenden Skizze ist M der Mittelpunkt des Quadrats und die Länge der Strecke von A nach B beträgt ein Viertel der Seitenlänge des Quadrats. Wie berechnet man den eingezeichneten Winkel?



★ Trigonometrie im Computer

In PYTHON bietet die Bibliothek MATH alle gängigen trigonometrischen Funktionen und die Konstante pi. Man beachte den [Rundungsfehler](#):

```
from math import *
sin(pi), cos(pi)
```

Wie die meisten Programmiersprachen erwartet PYTHON Angaben in [Bogenmaß](#). Obwohl die Umrechnung zwischen Grad und Radiant ganz einfach ist, gibt es dafür sogar fertige Funktionen, damit man sich das Denken abgewöhnen kann.

```
degrees(pi), radians(180)
```

Die Arkusfunktionen haben typische Kurznamen, z. B. asin statt arcsin.

```
asin(1/sqrt(2))/pi
```

Speziell für den Arkustangens gibt es eine sinnvolle zweistellige Erweiterung namens [atan2](#). Warum es diese Funktion in vielen Programmiersprachen gibt, sollte im Laufe von Abschnitt 19 klar geworden sein. Es wird zudem in dem Video [Der "bessere" Arkustangens](#) noch einmal ausführlich erklärt.

```
atan(3/4), atan2(3,4), atan2(-3,-4)
```

Für das Umrechnen zwischen kartesischen Koordinaten und der Polardarstellung gibt es Funktionen in der CMATH-Bibliothek.

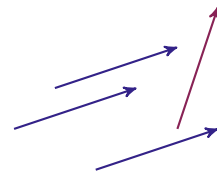
```
import cmath
cmath.polar(3+2j), cmath.rect(2,pi/4)/sqrt(2)
```

Übrigens kann man in PYTHON den [Operator für den Rest](#) auch für Fließkommazahlen verwenden. Das ist im Zusammenhang mit Winkeln manchmal praktisch:

```
42 % (2*pi)
```

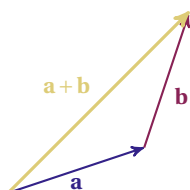
20. Vektoren

Vektoren sind von der Anschauung her so etwas wie Pfeile; sie haben eine Länge und eine Richtung. Sie haben allerdings keinen festen Ort! Die drei blauen Pfeile am Rand stellen z. B. alle *denselben* Vektor dar, während der rote Pfeil einen *anderen* Vektor repräsentiert. Vektoren kann man geometrisch als Darstellungen von Verschiebungen interpretieren. Sie werden u. a. in der Physik immer dann verwendet, wenn es darum geht, mit *gerichteten Größen* wie Geschwindigkeit, Beschleunigung oder Kraft zu arbeiten.



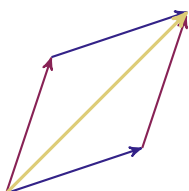
Wir werden für die Bezeichnung von Vektoren meistens fettgedruckte Kleinbuchstaben wie **a**, **v** oder **x** verwenden. Wenn man mit der Hand schreibt, macht man stattdessen einen Pfeil über den Buchstaben: $\vec{a}$. (Gebräuchlich ist mancherorts auch noch die Konvention, für Vektoren im Druck kleine **Frakturbuchstaben** zu verwenden. Das handschriftliche Pendant ist die **Sütterlinschrift**.)

Die *Addition* von zwei Vektoren wird anschaulich folgendermaßen definiert: Man verschiebt den Anfangspunkt des einen Vektors an den Endpunkt des anderen. Die Summe der beiden ist dann der Vektor, der vom Anfangspunkt des nicht verschobenen zum Endpunkt des verschobenen Vektors zeigt. Die beiden Vektoren von oben addiert man z. B. so:



Zur Erinnerung: Das entspricht der **Addition von komplexen Zahlen**.

Es stellt sich heraus, dass das Ergebnis unabhängig davon ist, welchen der beiden Summanden man verschiebt: die Vektoraddition ist **kommutativ**.

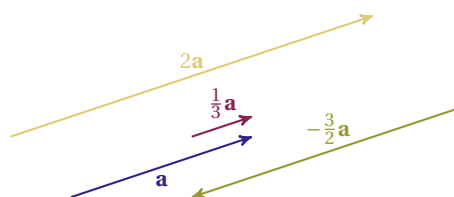


Das Bild, das sich ergibt, ist das sogenannte **Kräfteparallelogramm**, das Sie evtl. noch aus der Physik kennen. Man könnte auch sagen, dass die Vektoraddition so definiert ist, dass sie der Addition von Kräften in der realen Welt entspricht.

Man kann Vektoren außerdem mit einem Faktor multiplizieren. Dieser Faktor ist selbst kein Vektor, sondern eine Zahl und wird – zur Unterscheidung von den Vektoren – **Skalar** genannt. Die Multiplikation eines Skalars mit einem Vektor nennt man auch **Skalarmultiplikation**. Bei der Multiplikation mit einem Skalar wird ein Vektor um den entsprechenden Faktor verlängert bzw. verkürzt, ohne dass er seine Richtung ändert. (Er wird *skaliert*, daher auch das Wort.) Ist der Faktor negativ, so

Skalar
Skalarmultiplikation

ändert sich außerdem die Orientierung des Vektors. (Und Multiplikation mit 1 soll einfach wieder den Vektor selbst ergeben.)



Für Skalare werden wir meistens kleine griechische Buchstaben verwenden. Die Skalarmultiplikation schreibt man in der Form $\lambda \cdot \mathbf{a}$, lässt aber fast immer den Punkt weg: $\lambda \mathbf{a}$.

Wie man an der Skizze sofort sieht, sind Vektoren zu ihren skalaren Vielfachen parallel. Und das gilt auch umgekehrt: Wenn zwei Vektoren parallel sind, dann sind sie skalare Vielfache voneinander.

Nullvektor

Wir führen noch einen besonderen Vektor ein, der keine Richtung und keine Länge hat. Wir nennen ihn den *Nullvektor* und schreiben ihn als $\mathbf{0}$. Wenn man den Nullvektor zu einem beliebigen anderen Vektor $\mathbf{a}$ addiert, so kommt vereinbarungsgemäß wieder $\mathbf{a}$ heraus.²⁸ Der Nullvektor ist außerdem das Ergebnis, das man erhält, wenn man einen Vektor mit dem Skalar 0 multipliziert.

Vektoraddition und Skalarmultiplikation „passen zusammen“. Wenn $\mathbf{a}$ ein Vektor ist, dann gilt z. B. die Beziehung $2\mathbf{a} = \mathbf{a} + \mathbf{a}$, die wir auch von ganz normalen Zahlen kennen. (Überzeugen Sie sich durch eine Skizze!) Außerdem gilt $(-1) \cdot \mathbf{a} + \mathbf{a} = \mathbf{0}$,²⁹ weshalb wir in Zukunft für $(-1) \cdot \mathbf{a}$ einfach $-\mathbf{a}$ schreiben werden.

Der Vollständigkeit halber fassen wir die bereits erwähnten Eigenschaften und einige weitere zusammen.

Lemma 20.1

Für beliebige Vektoren $\mathbf{a}$, $\mathbf{b}$ und $\mathbf{c}$ sowie beliebige Skalare λ und μ gilt:

- (i) $\mathbf{a} + \mathbf{0} = \mathbf{a}$
- (ii) $\mathbf{a} + \mathbf{b} = \mathbf{b} + \mathbf{a}$
- (iii) $\mathbf{a} + (\mathbf{b} + \mathbf{c}) = (\mathbf{a} + \mathbf{b}) + \mathbf{c}$
- (iv) $0 \cdot \mathbf{a} = \mathbf{0}$
- (v) $1 \cdot \mathbf{a} = \mathbf{a}$
- (vi) $(-1) \cdot \mathbf{a} + \mathbf{a} = \mathbf{0}$
- (vii) $\lambda \cdot (\mathbf{a} + \mathbf{b}) = \lambda \cdot \mathbf{a} + \lambda \cdot \mathbf{b}$
- (viii) $(\lambda + \mu) \cdot \mathbf{a} = \lambda \cdot \mathbf{a} + \mu \cdot \mathbf{a}$
- (ix) $(\lambda \cdot \mu) \cdot \mathbf{a} = \lambda \cdot (\mu \cdot \mathbf{a})$

Beachten Sie, dass bei den letzten drei Punkten die Zeichen $\cdot$ und $+$ links und rechts vom Gleichheitszeichen evtl. unterschiedliche Bedeutung haben: $\cdot$ kann für

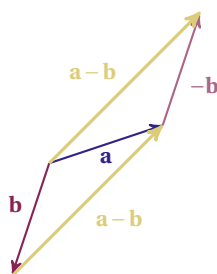
²⁸Der Nullvektor ist also das *neutrale Element* der Vektoraddition.

²⁹ $(-1) \cdot \mathbf{a}$ ist also bezüglich der Vektoraddition das *inverse Element* zu $\mathbf{a}$.

skalare Multiplikation stehen oder für die übliche Multiplikation reeller Zahlen. + kann für die Addition von Vektoren stehen oder für die von reellen Zahlen. Auch für die neuen Operationen gilt allerdings die übliche Konvention „Punkt- vor Strichrechnung“. Daher sind manche Ausdrücke geklammert und andere nicht.

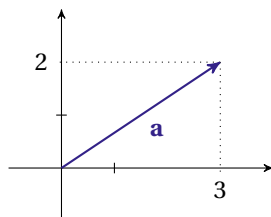
Man findet in manchen Büchern auch Ausdrücke wie $\mathbf{a}/\lambda$, wobei $\mathbf{a}$ ein Vektor und λ ein Skalar ist. Das ist keine neue Operation, sondern lediglich eine andere Schreibweise für die Skalarmultiplikation $1/\lambda \cdot \mathbf{a}$.

Welche geometrische Bedeutung hat die *Subtraktion* von Vektoren? Zunächst ist wie bei den Zahlen die Subtraktion eines Vektors lediglich die Addition des inversen Vektors: $\mathbf{a} - \mathbf{b}$ ist also nur eine andere Notation für $\mathbf{a} + (-\mathbf{b})$, wobei $-\mathbf{b}$ der Vektor ist, den man durch „Umdrehen“ aus $\mathbf{b}$ erhält. Eine Skizze macht deutlich, wie das zu interpretieren ist:



Wenn man $\mathbf{a}$ und $\mathbf{b}$ so verschiebt, dass sie beide denselben Anfangspunkt haben, dann ist $\mathbf{a} - \mathbf{b}$ der Vektor, der vom Endpunkt von $\mathbf{b}$ zum Endpunkt von $\mathbf{a}$ zeigt. (Erinnern Sie sich an Aufgabe 14.9.)

Wie Punkte, Geraden und Dreiecke sind Vektoren geometrische Objekte, die auch ohne Koordinatensystem „funktionieren“. Wenn man jedoch ein Koordinatensystem hat, dann kann man das Rechnen mit Vektoren auf das Rechnen mit Zahlen zurückführen. Man verschiebt einen Vektor so, dass sein Anfangspunkt im Ursprung des Koordinatensystems liegt und beschreibt ihn durch die Koordinaten seines Endpunktes.



Den Vektor in der obigen Skizze würde man also

$$\mathbf{a} = \begin{pmatrix} 3 \\ 2 \end{pmatrix} \quad (20.1)$$

schreiben. (Nach wie vor ist der Vektor aber *nicht* an einen festen Ort in der Ebene gebunden! Er wird lediglich zur Ermittlung seiner Koordinaten auf den Nullpunkt

verschoben.) Rein formal sind solche Ausdrücke **Tupel**, aber man schreibt die Komponenten häufig wie in (20.1) untereinander. Das verbraucht zwar mehr vertikalen Platz, wird sich aber noch als praktisch herausstellen.³⁰

Die Komponenten eines Vektors bezeichnet man manchmal dadurch, dass man an den „Namen“ des Vektors Indizes hängt (und dann auf Fettschrift bzw. Pfeile verzichtet):

$$\mathbf{b} = \begin{pmatrix} b_1 \\ b_2 \end{pmatrix}$$

Vektoraddition und Skalarmultiplikation sind nun ganz einfach zu berechnen, indem man *komponentenweise* vorgeht:

$$\begin{pmatrix} a_1 \\ a_2 \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} = \begin{pmatrix} a_1 + b_1 \\ a_2 + b_2 \end{pmatrix} \quad \lambda \cdot \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} = \begin{pmatrix} \lambda \cdot a_1 \\ \lambda \cdot a_2 \end{pmatrix} \quad (20.2)$$

(Machen Sie sich ggf. durch ein paar Skizzen klar, dass sich durch diese Art des Rechnens wirklich das ergibt, was wir vorher geometrisch durchgeführt haben.)

Wir werden im Folgenden zweigleisig verfahren. Einerseits interessieren uns Vektoren als Hilfsmittel für die Geometrie und damit rein technisch als Elemente der Mengen $\mathbb{R}^2$ (Ebene) bzw. $\mathbb{R}^3$ (Raum). Andererseits kann man rein formal Operationen wie in (20.2) mit jeder beliebigen Anzahl von Komponenten durchführen. Es ist egal, ob man in $\mathbb{R}^2$, $\mathbb{R}^3$ oder $\mathbb{R}^{42}$ rechnet und man kann sogar $\mathbb{R}$ durch $\mathbb{Q}$, $\mathbb{C}$ oder einen **Restklassenkörper** wie $\mathbb{Z}_7$ ersetzen. (Siehe z. B. Aufgabe 20.3.) Geometrisch kann man sich das zwar nicht vorstellen, aber man kann mit solchen Vektoren rechnen und in vielen Anwendungen wird das auch benötigt.

Aufgabe 20.1. Sei $\mathbf{a} = (-2, 5)$ und $\mathbf{b} = (1, -3)$. Berechnen Sie $\mathbf{a} + \mathbf{b}$ sowie $3 \cdot \mathbf{a}$ und $-1/2 \cdot \mathbf{b}$. Üben Sie mit ähnlichen, selbst gestellten Aufgaben, bis Sie solche Berechnungen sicher durchführen können.

Aufgabe 20.2. Rechnen Sie für jeden der Punkte aus Lemma 20.1 mindestens ein Beispiel mit selbstgewählten konkreten Werten durch, um sich von der Korrektheit zu überzeugen. (Und um sich klarzumachen, was genau eigentlich jeweils ausgesagt wird.) Wählen Sie Beispiele sowohl in $\mathbb{R}^2$ als auch in $\mathbb{R}^3$.

Aufgabe 20.3. Berechnen Sie für $\mathbf{a} = (3, 4)$, $\mathbf{b} = (2, 2)$ und $\lambda = 4$ die Vektoren $\mathbf{a} + \mathbf{b}$, $\mathbf{b} - \mathbf{a}$, $\lambda \cdot \mathbf{a}$ und $\lambda \cdot \mathbf{b}$. Dabei sollen jedoch sowohl die Komponenten der Vektoren als auch die Skalare Elemente von $\mathbb{Z}_5$ sein und alle Berechnungen sollen auch in diesem Restklassenkörper durchgeführt werden.

Wenn wir Vektoren als Tupel darstellen, dann mischen wir eigentlich zwei Sorten von Objekten, die man auseinanderhalten sollte: *Punkte* (die wir ebenfalls als Tupel darstellen) und *Vektoren*. Punkte haben einen festen Platz in der Ebene (oder im

³⁰Wegen des Platzverbrauchs werden Vektoren hier im Skript dann aber doch oft horizontal gesetzt, während ich sie an der Tafel eher vertikal schreiben werde. Zerbrechen Sie sich darüber nicht den Kopf. Die Unterscheidung ist mathematisch nicht relevant.

Raum) und bezeichnen einen Ort, während Vektoren beliebig verschiebbar sind und eine gerichtete Größe bezeichnen. Wir haben bereits einen gewissen Zusammenhang hergestellt, indem wir für die Darstellung von Vektoren die Koordinaten ihrer Endpunkte verwendet haben. Das kann man auch umdrehen: man kann einen Ort (Punkt) durch einen Vektor repräsentieren; und zwar durch den, der vom Ursprung des Koordinatensystems auf diesen Punkt zeigt. Das nennt man einen *Ortsvektor*. Der Ortsvektor des Punktes $(-2, 5)$ ist also z. B. der Vektor $\begin{pmatrix} -2 \\ 5 \end{pmatrix}$, für den wir auch manchmal einfach $(-2, 5)$ schreiben.

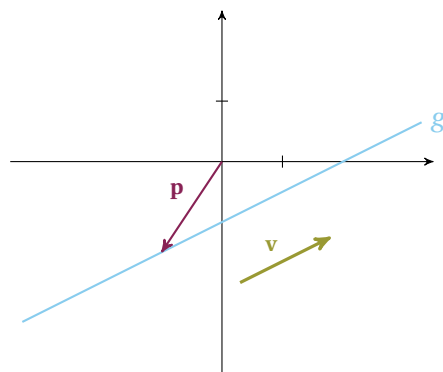
Ortsvektor

Man kann auf diese Ortsvektoren theoretisch die Rechenoperationen anwenden, die wir in diesem Kapitel eingeführt haben. Und man kann die beiden Vektortypen auch mischen. Es stellt sich aber heraus, dass es nur zwei Operationen gibt, die geometrisch einen Sinn ergeben:

- Man kann einen Ortsvektor (also einen Punkt) und einen normalen Vektor „addieren“. Das Ergebnis ist wieder ein Ortsvektor und so zu verstehen, dass der ursprüngliche Punkt um den Vektor verschoben wurde.
- Man kann zwei Ortsvektoren (also Punkte) „subtrahieren“. Das Ergebnis ist ein Vektor (*kein* Punkt!), der (siehe die Skizze zur Subtraktion) von einem Punkt zum anderen zeigt. Er gibt an, wie man den zweiten Punkt verschieben müsste, um zum ersten zu gelangen.

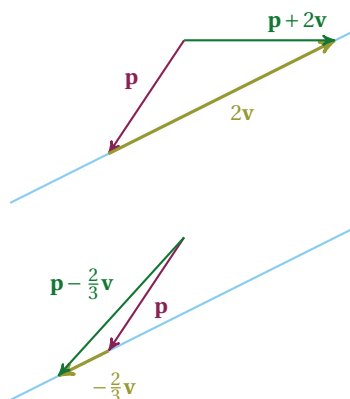
Da man in der Schreibweise keinen Unterschied zwischen normalen Vektoren und Ortsvektoren macht, muss man etwas aufpassen und sich jeweils klarmachen, welcher Vektor welche Rolle spielt.³¹

Als erste Anwendung der Konzepte aus diesem Abschnitt wiederholen wir die Darstellung von Geraden aus der Schule. Wie [am Anfang des Kapitels](#) bereits angekündigt, betrachten wir eine Gerade g als die *Menge* der Punkte, die auf g liegen. Um diese Menge zu beschreiben, wählen wir einen Punkt auf der Geraden und dazu den Ortsvektor $\mathbf{p}$. Außerdem noch einen Vektor $\mathbf{v}$, der parallel zur Geraden verläuft.



³¹Wenn Sie das verwirrend finden, stellen Sie sich stattdessen vor, dass es gar keine Ortsvektoren gibt. Der „Ortsvektor“ zum Punkt P ist dann einfach der „Schiebevektor“, der den Nullpunkt auf den Punkt P schiebt.

Vielfache des Vektors $\mathbf{v}$ können wir offenbar verwenden, um Punkte auf der Geraden zu verschieben. Insbesondere können wir vom Punkt $\mathbf{p}$ ausgehend³² alle anderen Punkte auf der Geraden erreichen.



Man kann g also auch so schreiben:

$$g = \{\mathbf{p} + \lambda \mathbf{v} : \lambda \in \mathbb{R}\}$$

Oder in diesem konkreten Beispiel:

$$g = \left\{ \begin{pmatrix} -1 \\ -3/2 \end{pmatrix} + \lambda \begin{pmatrix} 3/2 \\ 3/4 \end{pmatrix} : \lambda \in \mathbb{R} \right\} \quad (20.3)$$

Punkt-Richtungs-
Form
Richtungsvektor

Man nennt diese Darstellung eine *Punkt-Richtungs-Form* der Geraden und schreibt sie abkürzend auch als $\mathbf{p} + \mathbb{R}\mathbf{v}$. $\mathbf{v}$ ist ein *Richtungsvektor* der Geraden.

Beachten Sie, dass wir oben mit einem *beliebigen* Punkt $\mathbf{p}$ auf der Geraden angefangen haben und dass auch $\mathbf{v}$ *irgendein* Vektor parallel zu g war. Das impliziert, dass es für jede Gerade unendlich viele Möglichkeiten gibt, sie in Punkt-Richtungs-Form darzustellen. Insbesondere können zwei verschiedene Darstellungen dieselbe Gerade beschreiben.

Aufgabe 20.4. Setzen Sie verschiedene Werte für λ in Gleichung (20.3) ein und berechnen Sie jeweils die sich ergebenden Punkte. Kontrollieren Sie anhand der Skizze, dass die Punkte auf der Geraden liegen.

Aufgabe 20.5. Wandeln Sie die Darstellung (20.3) in eine *Geradengleichung* der Form $y = mx + b$ um, die dieselbe Gerade beschreibt.

Aufgabe 20.6. Setzen Sie die Punkte aus Aufgabe 20.4 zur Kontrolle in die Gleichung aus Aufgabe 20.5 ein.

Aufgabe 20.7. Welche Geraden kann man durch Gleichungen der Form $y = mx + b$ *nicht* beschreiben, durch die Punkt-Richtungs-Form aber schon?

³²Wir verwenden hier und in Zukunft die Konvention, dass wir vom *Punkt* $\mathbf{p}$ sprechen, obwohl eigentlich der *Ortsvektor* des Punktes gemeint ist. Wenn wir uns präziser ausdrücken würden, müssten wir viel mehr schreiben.

Aufgabe 20.8. Ermitteln Sie mithilfe von Gleichung (20.3), ob die Punkte $(2, 0)$ bzw. $(3, 1)$ auf der Geraden g liegen.

Aufgabe 20.9. Geben Sie mindestens eine andere Darstellung von g in Punkt-Richtungs-Form an.

Aufgabe 20.10. Die erste Komponente des allgemeinen Punktes (x, y) in der Darstellung (20.3) ist $x = -1 + \lambda \cdot 3/2$. Berechnen Sie entsprechend auch die zweite Komponente y . Lösen Sie dann den obigen Term für x nach λ auf und setzen Sie das Ergebnis in den Term für y ein. Was erhalten Sie?

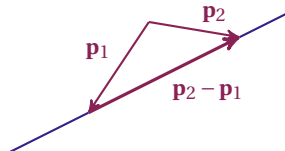
Aufgabe 20.11. Geben Sie die Gerade

$$\left\{ \begin{pmatrix} 3 \\ -2 \end{pmatrix} + \lambda \begin{pmatrix} 2 \\ -1 \end{pmatrix} : \lambda \in \mathbb{R} \right\}$$

in der Form $y = mx + b$ an und verwenden Sie dabei Aufgabe 20.10.

Aufgabe 20.12. Wenn man eine Gerade der Form $y = mx + b$ in Punkt-Richtungs-Form schreiben will, wie kommt man dann möglichst einfach an einen Punkt auf der Geraden und an einen Richtungsvektor?

Wenn zwei Punkte $\mathbf{p}_1$ und $\mathbf{p}_2$ gegeben sind, wird durch diese eindeutig eine Gerade bestimmt. In diesem Fall kann man die Differenz der beiden Ortsvektoren als Richtungsvektor wählen.



Die Geradengleichung sieht dann folgendermaßen aus

$$\{\mathbf{p}_1 + \lambda(\mathbf{p}_2 - \mathbf{p}_1) : \lambda \in \mathbb{R}\} = \{\lambda\mathbf{p}_2 + (1 - \lambda)\mathbf{p}_1 : \lambda \in \mathbb{R}\}$$

und die zweite Form hätte man auch so schreiben können:³³

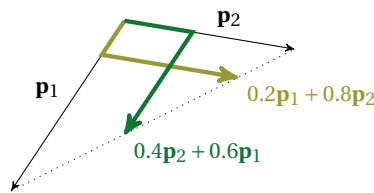
$$\{\lambda\mathbf{p}_1 + \mu\mathbf{p}_2 : \lambda, \mu \in \mathbb{R} \wedge \lambda + \mu = 1\} \quad (20.4)$$

Diese alternative Form der Geradendarstellung ist auch deshalb interessant, weil man aus ihr eine Darstellung der *Strecke* zwischen den beiden Punkten bekommt. Im Extremfall $\lambda = 1$ ist $\mu = 0$ und es gilt $\lambda\mathbf{p}_1 + \mu\mathbf{p}_2 = \mathbf{p}_1$. Ist umgekehrt $\mu = 1$, so liefert dieselbe Gleichung $\mathbf{p}_2$. Sind λ und μ Zahlen *zwischen* 0 und 1, so erhält man die Punkte *zwischen* $\mathbf{p}_1$ und $\mathbf{p}_2$.

³³Einen Ausdruck der Form $\lambda\mathbf{p}_1 + \mu\mathbf{p}_2$, wobei $\lambda + \mu = 1$ gilt, nennt man eine *Affinkombination* der Punkte $\mathbf{p}_1$ und $\mathbf{p}_2$. Lässt man die Forderung $\lambda + \mu = 1$ fallen, sind λ und μ also frei wählbar, so spricht man von einer *Linearkombination*. Letzterem Ausdruck sind Sie in einem anderen Zusammenhang schon in Kapitel II begegnet.

Affinkombination

Linearkombination



Die Menge der Punkte auf der Strecke von $\mathbf{p}_1$ nach $\mathbf{p}_2$ sieht demnach so aus:

$$\{\lambda \mathbf{p}_1 + \mu \mathbf{p}_2 : \lambda, \mu \in [0, 1] \wedge \lambda + \mu = 1\} \quad (20.5)$$

Aufgabe 20.13. Geben Sie für die beiden durch $y = -2x + 5$ bzw. $y = 12$ gegebenen Geraden jeweils eine Punkt-Richtungs-Form an.

Aufgabe 20.14. Geben Sie eine Punkt-Richtungs-Form der Geraden durch die Punkte $P_1 = (1, 2)$ und $P_2 = (-3, 1)$ an.

Aufgabe 20.15. Liegt $(5, 3)$ auf der Geraden aus Aufgabe 20.14? Und liegt der Punkt auf der Strecke zwischen P_1 und P_2 ?

Aufgabe 20.16. Seien $\mathbf{a}$ und $\mathbf{b}$ zwei verschiedene Punkte in der Ebene, die beide nicht der Nullpunkt sind. Welche der Punkte

$$\mathbf{a} + \mathbf{b}, \quad \mathbf{a} - \mathbf{b}, \quad \frac{1}{3} \cdot (2\mathbf{a} + \mathbf{b}), \quad \frac{1}{3} \cdot (\mathbf{a} + 2\mathbf{b}) \quad \text{und} \quad \frac{1}{3} \cdot (4\mathbf{a} - \mathbf{b})$$

liegen auf der Verbindungsgeraden von $\mathbf{a}$ und $\mathbf{b}$? (Wie üblich werden hier Punkte mit ihren Ortsvektoren identifiziert.)

Aufgabe 20.17. Wie würde sich die Antwort auf Aufgabe 20.16 ändern, wenn man das Wort *Verbindungsgerade* durch das Wort *Verbindungsstrecke* ersetzen würde?

Aufgabe 20.18. Warum wurde in Aufgabe 20.16 der Fall ausgeschlossen, dass einer der beiden Punkte der Nullpunkt ist?

Aufgabe 20.19. Seien $\mathbf{p}$ und $\mathbf{v}$ die Vektoren $(4, 5, 6)$ und $(1, -2, 3)$ und g die Gerade $\mathbf{p} + \mathbb{R}\mathbf{v}$. Welchen Wert muss a haben, damit der Punkt $(1, a, -3)$ auf g liegt?

Betrachtet man geometrische Figuren als Mengen von Punkten, so besteht der **Durchschnitt** zweier solcher Mengen aus der Menge aller Punkte, die zu beiden Figuren gehören. Das passt auch zu der üblichen Formulierung, dass sich beispielsweise eine Gerade und ein Kreis *schneiden*. Wie sieht beispielsweise die Schnittmenge der folgenden beiden Geraden in Punkt-Richtungs-Form aus?

$$g_1 = \begin{pmatrix} 3 \\ -1 \end{pmatrix} + \mathbb{R} \begin{pmatrix} 2 \\ 7 \end{pmatrix} \quad g_2 = \begin{pmatrix} 2 \\ 2 \end{pmatrix} + \mathbb{R} \begin{pmatrix} -1 \\ 5 \end{pmatrix}$$

Jeder Punkt $\mathbf{p}$, der auf *beiden* Geraden liegt, muss sich sowohl als

$$\mathbf{p} = \begin{pmatrix} 3 \\ -1 \end{pmatrix} + \lambda \begin{pmatrix} 2 \\ 7 \end{pmatrix} \quad \text{als auch als} \quad \mathbf{p} = \begin{pmatrix} 2 \\ 2 \end{pmatrix} + \mu \begin{pmatrix} -1 \\ 5 \end{pmatrix}$$

darstellen lassen.³⁴ Man kann das also als

$$\begin{pmatrix} 3 \\ -1 \end{pmatrix} + \lambda \begin{pmatrix} 2 \\ 7 \end{pmatrix} = \begin{pmatrix} 2 \\ 2 \end{pmatrix} + \mu \begin{pmatrix} -1 \\ 5 \end{pmatrix}$$

bzw. komponentenweise als

$$\begin{aligned} 3 + 2\lambda &= 2 - \mu \\ -1 + 7\lambda &= 2 + 5\mu \end{aligned}$$

in Form von zwei Gleichungen schreiben. Etwas umsortiert so:

$$\begin{aligned} 2\lambda + \mu &= -1 \\ 7\lambda - 5\mu &= 3 \end{aligned}$$

Sie sollten sich aus der Schule daran erinnern, dass wir es hier mit einem *linearen Gleichungssystem* (LGS) zu tun haben. Und Sie sollten – zumindest für kleine Systeme aus zwei oder drei Gleichungen – auch wissen, wie man so etwas löst. Wir werden uns aber demnächst noch ausführlicher mit linearen Gleichungssystemen beschäftigen. Momentan reicht es zu wissen, dass man es im Zusammenhang mit geometrischen Fragen mit solchen Systemen zu tun bekommt.

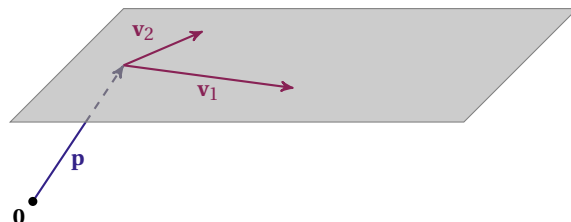
Aufgabe 20.20. Berechnen Sie den Schnittpunkt von g_1 und g_2 , d. h. den Punkt $\mathbf{p}$ mit $g_1 \cap g_2 = \{\mathbf{p}\}$.

Aufgabe 20.21. Im obigen Fall hatten die beiden Geraden genau einen gemeinsamen Punkt. Welche Fälle können noch eintreten, wenn man zwei Geraden in der Ebene hat? Und wie ist das im Raum?

Zum Schluss sei noch darauf hingewiesen, dass man auch im Raum Geraden durch eine Punkt-Richtungs-Form darstellen kann. Es ändert sich eigentlich gar nichts, außer dass man es nun mit drei statt zwei Komponenten zu tun hat.

Aufgabe 20.22. Geben Sie eine Punkt-Richtungs-Form der Geraden im Raum durch die Punkte $P_1 = (1, 2, 0)$ und $P_2 = (-3, 1, 5)$ an.

Im Raum kann man aber noch etwas anderes machen. Man kann eine *Ebene* eindeutig dadurch beschreiben, dass man einen Ortsvektor und *zwei* Richtungsvektoren angibt.



³⁴Man beachte, dass wir hier zwei verschiedene Variablennamen, λ und μ , gewählt haben. Es gibt keinen Grund anzunehmen, dass der Parameter in beiden Darstellungen identisch sein muss.

Man erhält also eine Punkt-Richtungs-Form für Ebenen:

$$\{\mathbf{p} + \lambda \mathbf{v}_1 + \mu \mathbf{v}_2 : \lambda, \mu \in \mathbb{R}\} = \mathbf{p} + \mathbb{R} \mathbf{v}_1 + \mathbb{R} \mathbf{v}_2 \quad (20.6)$$

Aufgabe 20.23. Welche Bedingung müssen die beiden Richtungsvektoren $\mathbf{v}_1$ und $\mathbf{v}_2$ erfüllen, damit durch (20.6) wirklich eine Ebene beschrieben wird?

Aufgabe 20.24. Geben Sie eine Punkt-Richtungs-Form der Ebene an, die die drei Punkte $P_1 = (1, 2, 1)$, $P_2 = (3, -1, 5)$ und $P_3 = (4, 0, -6)$ enthält.

Aufgabe 20.25. Was für ein Gleichungssystem erhält man, wenn man die Schnittmenge der Geraden aus Aufgabe 20.22 und der Ebene aus Aufgabe 20.24 ermitteln will? (Sie müssen den Schnitt nicht berechnen. Stellen Sie nur die Gleichungen auf.)

Aufgabe 20.26. Was kann ganz allgemein herauskommen, wenn man die Schnittmenge einer Geraden und einer Ebene berechnet? Und welche Fälle können eintreten, wenn man zwei Ebenen schneidet?

★ Vektoren im Computer

Es gibt natürlich diverse Bibliotheken, die mit Vektoren arbeiten können. Die einfachen Rechenoperationen aus diesem Abschnitt kann man jedoch auch selbst implementieren. Stellt man in PYTHON Vektoren als Listen dar, so kann man es beispielsweise so machen:

```
# Addition von Vektoren
vectorAdd = lambda V,W: [v+w for v,w in zip(V,W)]
# skalare Multiplikation
scalarVectorMult = lambda s,V: [s*v for v in V]
```

VI

Lineare Algebra

Obwohl es zunächst nicht den Anschein hat, wird es auch in den kommenden Abschnitten weiterhin um Geometrie gehen. Aber wir verwenden nun vergleichsweise moderne Methoden,¹ die auch außerhalb der Geometrie in vielen Bereichen Anwendung finden. Im Prinzip haben wir damit schon angefangen, als wir **Vektoren** in Koordinatenschreibweise notiert haben. „Geometrie mit Koordinaten“ nennt man *analytische Geometrie*. Die gibt es erst seit ungefähr vier Jahrhunderten. Die seit der Antike betriebene „Geometrie ohne Koordinaten“ nennt man heutzutage *synthetische Geometrie*.

Achtung: Schon im **Abschnitt über Vektoren** haben wir damit angefangen (z. B. Aufgabe 20.3), statt in $\mathbb{R}$ in **Restklassenkörpern** zu rechnen. Das werden wir auch im Folgenden ab und zu machen. Die grundsätzliche Regel ist, dass alle Operationen, für die nur die Grundrechenarten (also keine Wurzeln, trigonometrische Funktionen oder ähnliche Dinge) benötigt werden, in Restklassenkörpern genauso funktionieren wie in den reellen Zahlen. Auch wenn es in der Regel nicht explizit gesagt wird, können Sie also an so ziemlich allen Stellen $\mathbb{R}$ durchgehend durch einen Restklassenkörper der Form $\mathbb{Z}_p$ (oder auch durch $\mathbb{Q}$ oder $\mathbb{C}$) ersetzen.

21. Matrizen

In diesem Abschnitt werden wir uns mit *Matrizen* beschäftigen. Wir werden sie u. a. beim Lösen von **linearen Gleichungssystemen** und bei der Beschreibung von **linearen Abbildungen** benutzen. Sie werden aber auch in vielen anderen Bereichen der Mathematik und in anderen Wissenschaften als universelles Werkzeug eingesetzt.

Eine *Matrix* ist zunächst einmal nichts weiter als ein rechteckiges Schema von Zahlen (oder noch allgemeiner von irgendwelchen mathematischen Objekten). In der Regel umschließt man dieses Schema mit runden Klammern.

Matrix

$$\mathbf{A} = \begin{pmatrix} 3 & 4 & -2 & 0 \\ 1/2 & \pi & 42 & -\sqrt{2} \\ -7 & 1 & 1 & -2/7 \end{pmatrix} \quad (21.1)$$

¹In meinem Video *Was ist lineare Algebra?* sage ich etwas zum geschichtlichen Hintergrund.

Komponente Für Matrizen werden wir fettgedruckte Großbuchstaben verwenden. Die einzelnen Einträge nennt man *Komponenten*. Man verwendet für die Komponenten eine ähnliche Bezeichnungsweise wie bei Vektoren, allerdings mit Doppelindizes, die auf Zeile und Spalte der Komponente verweisen.

$$\mathbf{A} = \begin{pmatrix} a_{1,1} & a_{1,2} & a_{1,3} & a_{1,4} \\ a_{2,1} & a_{2,2} & a_{2,3} & a_{2,4} \\ a_{3,1} & a_{3,2} & a_{3,3} & a_{3,4} \end{pmatrix}$$

Hauptdiagonale Die *Hauptdiagonale* einer Matrix besteht aus den Komponenten $a_{i,j}$, bei denen $i = j$ gilt. Diese sind hier grün markiert.

Typ Matrizen können offenbar unterschiedliche Größe bzw. Gestalt haben; man spricht auch vom *Typ* der Matrix. Eine Matrix mit m Zeilen und n Spalten bezeichnet man als $m \times n$ -Matrix. Die Matrix $\mathbf{A}$ aus (21.1) ist also eine 3×4 -Matrix. Sind Zeilen- und Spaltenzahl gleich, spricht man auch von einer *quadratischen Matrix*. Die Menge *aller* $m \times n$ -Matrizen, deren Komponenten alle reelle Zahlen sind, wird üblicherweise $\mathbb{R}^{m \times n}$ geschrieben. Für unser Beispiel gilt also $\mathbf{A} \in \mathbb{R}^{3 \times 4}$. Wie Sie sich aufgrund dieser Schreibweise schon denken können, kann man Matrizen auch über anderen Grundmengen definieren. Auch so etwas wie $\mathbb{Q}^{2 \times 3}$ oder $\mathbb{Z}_2^{5 \times 5}$ ergibt also Sinn.

Zunächst ohne Motivation definieren wir nun ein paar Rechenoperationen. Es wird sich später noch herausstellen, wozu man diese gebrauchen kann. Matrizen vom gleichen Typ werden addiert, indem man (wie bei Vektoren) *komponentenweise* addiert. Die Summe zweier Matrizen ist also wieder eine Matrix gleichen Typs. Zum Beispiel:

$$\begin{pmatrix} 1 & 2 \\ 0 & 8 \\ -1/2 & 20 \end{pmatrix} + \begin{pmatrix} 1 & -1 \\ \sqrt{3} & 7 \\ -1/2 & 22 \end{pmatrix} = \begin{pmatrix} 1+1 & 2-1 \\ 0+\sqrt{3} & 8+7 \\ -1/2-1/2 & 20+22 \end{pmatrix} = \begin{pmatrix} 2 & 1 \\ \sqrt{3} & 15 \\ -1 & 42 \end{pmatrix}$$

Auch wie bei Vektoren kann man mit einem Skalar multiplizieren:

$$(-3) \cdot \begin{pmatrix} 1 & -2 \\ \sqrt{2} & 0 \\ 2/3 & -14 \end{pmatrix} = \begin{pmatrix} (-3) \cdot 1 & (-3) \cdot (-2) \\ (-3) \cdot \sqrt{2} & (-3) \cdot 0 \\ (-3) \cdot 2/3 & (-3) \cdot (-14) \end{pmatrix} = \begin{pmatrix} -3 & 6 \\ -3\sqrt{2} & 0 \\ -2 & 42 \end{pmatrix}$$

Nullmatrix Eine Matrix, deren Einträge alle null sind, nennt man *Nullmatrix* und schreibt dafür oft $\mathbf{0}$ (wie für den Nullvektor). Es gibt für jeden Typ eine andere Nullmatrix, es sollte jedoch im Allgemeinen aus dem Kontext hervorgehen, welche Nullmatrix gemeint ist.

Da diese beiden Rechenoperationen für Matrizen genau wie die für Vektoren komponentenweise arbeiten, erhalten wir auch dieselben Eigenschaften, die wir für Vektoren schon kennen.

Sind $\mathbf{A}$, $\mathbf{B}$ und $\mathbf{C}$ beliebige Matrizen sowie λ und μ beliebige Skalare, so gilt:

- $\mathbf{A} + \mathbf{0} = \mathbf{A}$
- $\mathbf{A} + \mathbf{B} = \mathbf{B} + \mathbf{A}$
- $\mathbf{A} + (\mathbf{B} + \mathbf{C}) = (\mathbf{A} + \mathbf{B}) + \mathbf{C}$
- $\mathbf{0} \cdot \mathbf{A} = \mathbf{0}$
- $\mathbf{1} \cdot \mathbf{A} = \mathbf{A}$
- $(-1) \cdot \mathbf{A} + \mathbf{A} = \mathbf{0}$
- $\lambda \cdot (\mathbf{A} + \mathbf{B}) = \lambda \cdot \mathbf{A} + \lambda \cdot \mathbf{B}$
- $(\lambda + \mu) \cdot \mathbf{A} = \lambda \cdot \mathbf{A} + \mu \cdot \mathbf{A}$
- $(\lambda \cdot \mu) \cdot \mathbf{A} = \lambda \cdot (\mu \cdot \mathbf{A})$

Dabei wird natürlich vorausgesetzt, dass die Operationen definiert sind, d. h., dass die beteiligten Matrizen vom Typ her passen.

Lemma 21.1

Es ergibt wie bei den Vektoren Sinn, die Matrix $(-1) \cdot \mathbf{A}$ kurz einfach $-\mathbf{A}$ zu schreiben und $\mathbf{A} - \mathbf{B}$ als Abkürzung für $\mathbf{A} + (-\mathbf{B})$ zu verwenden.

Aufgabe 21.1. Wie üblich sollten Sie die Eigenschaften aus Lemma 21.1 alle einmal durchgehen und mit Beispieldaten überprüfen.

Auf den ersten Blick unnötig kompliziert scheint die Matrizenmultiplikation zu sein. Aus Gründen, die **später** klar werden, multipliziert man *nicht* einfach komponentenweise. Der Ausdruck $\mathbf{A} \cdot \mathbf{B}$ ist überhaupt nur definiert, wenn die Anzahl der Spalten von $\mathbf{A}$ der Anzahl der Zeilen von $\mathbf{B}$ entspricht. Ist $\mathbf{A}$ eine $m \times n$ -Matrix und $\mathbf{B}$ eine $n \times p$ -Matrix, so ist $\mathbf{A} \cdot \mathbf{B}$ eine $m \times p$ -Matrix und der Eintrag in der i -ten Zeile und k -ten Spalte des Produktes wird folgendermaßen berechnet:

$$\sum_{j=1}^n a_{i,j} b_{j,k} = a_{i,1} b_{1,k} + a_{i,2} b_{2,k} + \cdots + a_{i,n} b_{n,k}$$

Sie sollten sich diese Formel allerdings nicht merken, sondern stattdessen eine grafische Vorstellung entwickeln, die man in Deutschland *Falksches Schema* nennt:

Falksches Schema

$$\begin{pmatrix} 3 & 7 & 0 \\ -1 & -2 & 9 \\ -1 & 4 & 3 \\ 8 & 0 & -6 \end{pmatrix} \cdot \begin{pmatrix} 5 & 2 & -4 \\ -1 & 3 & 7 \\ 0 & -2 & 2 \end{pmatrix} = \begin{pmatrix} 8 & 4 & 8 \\ \vdots & \vdots & \vdots \\ \vdots & \vdots & \vdots \\ \vdots & \vdots & \vdots \end{pmatrix}$$

$$(-1) \cdot 2 + 4 \cdot 3 + 3 \cdot (-2) = 4$$

$$(-1) \cdot (-4) + (-2) \cdot 7 + 9 \cdot 2 = 8$$

Man stellt sich die beiden Matrizen so nebeneinander vor, dass der rechte Faktor nach oben versetzt wird, so dass man unter ihm Platz für das Produkt hat. Multiplizieren kann man nur dann, wenn die Spaltenzahl links und die Zeilenzahl oben (im Beispiel jeweils drei) übereinstimmen, wenn man also oben links quasi ein Quadrat hat.

Für jeden Eintrag des Produktes betrachtet man nun die auf ihn zeigende Spalte des oberen Faktors sowie die auf ihn zeigende Zeile des linken Faktors. Die entsprechenden Einträge werden paarweise miteinander multipliziert und die Produkte werden am Ende aufsummiert.

Aufgabe 21.2. Im obigen Beispiel wurden zwei Komponenten des Produktes bereits berechnet. Rechnen Sie die anderen zehn Einträge aus.

Aufgabe 21.3. Seien $\mathbf{A} \in \mathbb{R}^{3 \times 5}$, $\mathbf{B} \in \mathbb{R}^{3 \times 3}$ und $\mathbf{C} \in \mathbb{R}^{5 \times 3}$. Welche der folgenden Produkte kann man bilden und welchen Typ hat das Produkt dann jeweils?

$$\mathbf{A} \cdot \mathbf{A}, \quad \mathbf{A} \cdot \mathbf{B}, \quad \mathbf{A} \cdot \mathbf{C}, \quad \mathbf{B} \cdot \mathbf{A}, \quad \mathbf{B} \cdot \mathbf{B}, \quad \mathbf{B} \cdot \mathbf{C}, \quad \mathbf{C} \cdot \mathbf{A}, \quad \mathbf{C} \cdot \mathbf{B}, \quad \mathbf{C} \cdot \mathbf{C}$$

Aufgabe 21.4. Berechnen Sie $\mathbf{A} \cdot \mathbf{B}$ und $\mathbf{B} \cdot \mathbf{A}$ für die beiden folgenden Matrizen.

$$\mathbf{A} = \begin{pmatrix} 2 & 0 \\ 1 & 1 \end{pmatrix} \quad \mathbf{B} = \begin{pmatrix} 1 & 2 \\ 0 & 0 \end{pmatrix}$$

Aufgabe 21.5. Ist die Matrizenmultiplikation **kommutativ**, gilt also $\mathbf{A} \cdot \mathbf{B} = \mathbf{B} \cdot \mathbf{A}$ für *alle* (quadratischen) Matrizen $\mathbf{A}$ und $\mathbf{B}$?

Aufgabe 21.6. Berechnen Sie das folgende Produkt im Restklassenkörper $\mathbb{Z}_7$:

$$\begin{pmatrix} 5 & 2 & 6 \\ 5 & 2 & 3 \\ 5 & 6 & 2 \end{pmatrix} \cdot \begin{pmatrix} 3 & 0 & 0 \\ 4 & 2 & 1 \\ 4 & 6 & 0 \end{pmatrix}$$

Aufgabe 21.7. Wie müsste eine 2×2 -Matrix $\mathbf{X}$ aussehen, die für *jede* 2×2 -Matrix $\mathbf{A}$ die Gleichung $\mathbf{X} \cdot \mathbf{A} = \mathbf{A}$ erfüllt?

Einheitsmatrix

$\mathbf{E}_n$

Eine quadratische Matrix, in deren Hauptdiagonale nur Einsen stehen und deren sämtliche andere Einträge verschwinden, nennt man *Einheitsmatrix* (engl. *identity matrix*). Hat so eine Matrix den Typ $n \times n$, so verwendet man für sie eins der Symbole $\mathbf{E}_n$ (in diesem Skript) oder $\mathbb{1}_n$ (in manchen anderen Büchern). $\mathbf{E}_4$ sieht zum Beispiel so aus:

$$\mathbf{E}_4 = \mathbb{1}_4 = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

Eine Einheitsmatrix $\mathbf{E}$ hat die Eigenschaft, dass $\mathbf{A} \cdot \mathbf{E} = \mathbf{A}$ bzw. $\mathbf{E} \cdot \mathbf{B} = \mathbf{B}$ für alle Matrizen $\mathbf{A}$ bzw. $\mathbf{B}$ gilt, die man (von links oder rechts) mit $\mathbf{E}$ multiplizieren kann.

Einheitsmatrizen sind also *neutrale Elemente* bzgl. der Matrizenmultiplikation, d. h., sie verhalten sich wie die Zahl eins bei der normalen Multiplikation von Zahlen.

Diese und weitere Eigenschaften der Matrizenmultiplikation sind im folgenden Lemma zusammengefasst.

Sind $\mathbf{A}$, $\mathbf{B}$ und $\mathbf{C}$ beliebige Matrizen, so gelten die folgenden Aussagen:

- $\mathbf{A} \cdot (\mathbf{B} \cdot \mathbf{C}) = (\mathbf{A} \cdot \mathbf{B}) \cdot \mathbf{C}$
- $\mathbf{0} \cdot \mathbf{A} = \mathbf{0}$
- $\mathbf{A} \cdot \mathbf{0} = \mathbf{0}$
- $\mathbf{A} \cdot \mathbf{E}_n = \mathbf{A}$
- $\mathbf{E}_n \cdot \mathbf{A} = \mathbf{A}$
- $\mathbf{A} \cdot (\mathbf{B} + \mathbf{C}) = \mathbf{A} \cdot \mathbf{B} + \mathbf{A} \cdot \mathbf{C}$
- $(\mathbf{A} + \mathbf{B}) \cdot \mathbf{C} = \mathbf{A} \cdot \mathbf{C} + \mathbf{B} \cdot \mathbf{C}$

Dabei ist das natürlich immer so gemeint, dass die entsprechen Typen „passen“, so dass die jeweiligen Operationen (Addition oder Multiplikation) überhaupt durchgeführt werden können.

Lemma 21.2

Von der Korrektheit dieser Aussagen kann man sich durch simples Nachrechnen überzeugen. Wie üblich ist es empfehlenswert, das an Beispielen wirklich einmal durchzuprobieren, um Routine im Umgang mit Matrizen zu bekommen. Beachten Sie, dass im Lemma mehrere Eigenschaften scheinbar doppelt aufgeführt sind. Das liegt daran, dass die Matrizenmultiplikation *nicht kommutativ* ist.

Aufgabe 21.8. Sei $\mathbf{M}$ die folgende Matrix. Geben Sie den ersten Eintrag der ersten Zeile von $\mathbf{M}^4$ an. Dabei ist ein Ausdruck wie $\mathbf{M}^4$ natürlich wie üblich als Abkürzung für das Produkt $\mathbf{M} \cdot \mathbf{M} \cdot \mathbf{M} \cdot \mathbf{M}$ gemeint.

$$\mathbf{M} = \begin{pmatrix} a & 0 & 0 \\ 1 & 1/2 & 3 \\ \pi & 9 & -42 \end{pmatrix}$$

Hinweis: Beachten Sie die beiden Nullen in der ersten Zeile. Muss man alle Einträge von $\mathbf{M}^2$ kennen, um die Aufgabe zu lösen?

Aufgabe 21.9. Geben Sie eine 2×2 -Matrix $\mathbf{A}$ an, für die $\mathbf{A} \cdot \mathbf{A} = \mathbf{A}$ gilt. Die Nullmatrix und die Einheitsmatrix zählen *nicht* als Lösungen. (Hinweis: Probieren Sie ein paar einfache Matrizen durch. Da zwei Matrizen explizit ausgeschlossen wurden, könnte man es mit einfachen Variationen derselben versuchen.)

Nun führen wir ein, was es bedeuten soll, dass eine Matrix mit einem Vektor multipliziert wird. Dabei stellt man sich vor, dass ein Vektor eine Matrix mit nur einer Spalte und entsprechend vielen Zeilen ist. Dann wird die normale Matrizenmultiplikation durchgeführt. Da bei dieser Art der Multiplikation per Definition der Vektor immer der rechte Faktor ist, ist das Ergebnis wieder eine einspaltige Matrix, also ein Vektor.

$$\begin{array}{c}
 \xrightarrow{\quad} \downarrow \begin{pmatrix} 2 \\ 3 \\ -2 \end{pmatrix} \\
 \begin{pmatrix} 4 & 5 & 0 \\ -2 & -2 & 8 \\ -1 & 4 & 3 \\ 7 & 0 & -6 \end{pmatrix} \begin{pmatrix} 4 \end{pmatrix}
 \end{array}$$

Kurz und knackig kann man das als „Matrix mal Vektor ist Vektor“ zusammenfassen. Es sollte damit auch klar sein, wieso man Vektoren oft vertikal schreibt: Im Zusammenspiel mit Matrizen betrachtet man sie als „schlanke“ Matrizen. Man beachte aber, dass diese Multiplikation nur dann funktioniert, wenn die Typen passen: Die Anzahl der Komponenten des ersten Vektors muss der Anzahl der Spalten der Matrix entsprechen und die Anzahl der Komponenten des zweiten Vektors wird die Anzahl der Zeilen der Matrix sein.

Aufgabe 21.10. Berechnen Sie im obigen Beispiel die fehlenden Komponenten des Produktvektors.

Transposition

Manchmal ist es auch hilfreich, Matrizen zu *transponieren*. Damit ist gemeint, dass die Matrix an ihrer Hauptdiagonalen gespiegelt wird. Man könnte es auch so ausdrücken, dass Zeilen und Spalten vertauscht werden: aus der ersten Zeile wird die erste Spalte, aus der zweiten Zeile die zweite Spalte, etc. (Eine Matrix, die nicht quadratisch ist, ändert dabei auch ihren Typ!) Für die Transponierte einer Matrix $\mathbf{A}$ schreibt man $\mathbf{A}^T$.

$$\mathbf{A} = \begin{pmatrix} 4 & 5 & 0 \\ -2 & -2 & 8 \\ -1 & 4 & 3 \\ 7 & 0 & -6 \end{pmatrix} \quad \mathbf{A}^T = \begin{pmatrix} 4 & -2 & -1 & 7 \\ 5 & -2 & 4 & 0 \\ 0 & 8 & 3 & -6 \end{pmatrix}$$

Beachten Sie, dass im obigen Beispiel die rot markierten Komponenten der Hauptdiagonale ihren Platz beibehalten haben. Außerdem ist aus einer 4×3 - eine 3×4 -Matrix geworden.

Aufgabe 21.11. Ergänzen Sie in der folgenden Matrix $\mathbf{A}$ die Einträge a , b und c so, dass $\mathbf{A}^T = \mathbf{A}$ gilt.

$$\mathbf{A} = \begin{pmatrix} \frac{1}{2} & -3 & a & -3 \\ -3 & 42 & -3 & 7 \\ 1 & -3 & 2\pi & b \\ -3 & c & 2 & -\sqrt{8} \end{pmatrix}$$

Aufgabe 21.12. Für die folgende Matrix gilt $\mathbf{M}^T = \mathbf{M}$. Geben Sie d an.

$$\mathbf{M} = \begin{pmatrix} a & 2 & b \\ a & b & a-2b \\ 2a & d & c \end{pmatrix}$$

In diesem Sinne kann man auch Vektoren transponieren: Wird ein Vektor $\mathbf{v}$ mit n Komponenten wie oben als $n \times 1$ -Matrix behandelt, dann ist mit $\mathbf{v}^T$ die „liegende“ $1 \times n$ -Matrix gemeint.

Es gibt einen wichtigen Zusammenhang zwischen Transposition und Matrizenmultiplikation, den wir noch öfter verwenden werden. Sind nämlich $\mathbf{A}$ und $\mathbf{B}$ Matrizen, die man multiplizieren kann, so gilt

$$(\mathbf{A} \cdot \mathbf{B})^T = \mathbf{B}^T \cdot \mathbf{A}^T, \quad (21.2)$$

d. h., das Transponieren dreht quasi die Reihenfolge der Multiplikation um.

Aufgabe 21.13. Zeichnen Sie schematische Skizzen nach dem [Falkschen Schema](#), um sich zu vergewissern, dass (21.2) korrekt ist.

Aufgabe 21.14. Machen Sie sich klar, dass $(\mathbf{A} \cdot \mathbf{B})^T = \mathbf{A}^T \cdot \mathbf{B}^T$ – also (21.2) „ohne Vertauschen“ – allein schon deshalb nicht stimmen kann, weil dann die Typen der Faktoren auf der rechten Seite des Gleichheitszeichens nicht stimmen würden.

Aufgabe 21.15. Was bekommt man, wenn man einen transponierten Vektor von links mit einer (passenden) Matrix multipliziert?

Aufgabe 21.16. Geben Sie eine 2×2 -Matrix $\mathbf{A}$ an, für die $\mathbf{A} \cdot \mathbf{A} = \mathbf{A}^T$ gilt. Die Nullmatrix und die Einheitsmatrix zählen *nicht* als Lösungen.

★ Matrizen im Computer

Stellt man in PYTHON Matrizen als Listen von Zeilen dar, die selbst Listen sind, so kann man mithilfe von [Funktion aus dem letzten Kapitel](#) die folgenden einfachen Funktionen für das Rechnen mit Matrizen definieren:

```
# Addition von Matrizen
matrixAdd = lambda A,B: [vectorAdd(a,b) for a,b in zip(A, B)]
# skalare Multiplikation von Matrizen
scalarMatrixMult = lambda x,A: [scalarVectorMult(x,a) for a in A]
# Hilfsfunktion für das komponentenweise Multiplizieren
dotProd = lambda A,B: sum(a*b for a,b in zip(A,B))
# i-te Spalte der Matrix A
col = lambda A,i: [row[i] for row in A]
# Produkt zweier Matrizen
def matrixProd(A,B):
    return [[scalarProd(row, col(B,i))
              for i in range(len(B[0]))] for row in A]
# Matrix mal Vektor, damit man V nicht als Matrix eingeben muss
def matTimesVec (A,V):
    return [p[0] for p in matrixProd(A, [[v] for v in V])]
# Transposition
transpose = lambda A: [col(A, i) for i in range(len(A[0]))]
```

22. Lineare Gleichungssysteme

Lineare Gleichungssysteme kamen bereits im [Abschnitt über Vektoren](#) vor. Sie spielen sowohl in der Geometrie (und damit auch in der Computergrafik) als auch in vielen anderen Bereichen der Mathematik und in Anwendungen eine zentrale Rolle. Und zwar in gewissem Sinne schlicht und einfach deshalb, weil es die einzigen Gleichungssysteme sind, für die es ein allgemeines Lösungsverfahren gibt. Daher versucht man oft, Probleme so zu formulieren, dass man sie mithilfe von linearen Gleichungssystemen lösen kann.

Gleichungen kennt man schon aus der Schule. Eine ganz simple wäre etwa $x^2 - 2 = 23$. Man sagt, dass man so eine Gleichung *löst* (oder präziser: *nach x auflöst*), wenn man einen oder mehrere Werte findet, die man für x einsetzen kann, so dass aus der Gleichung eine wahre Aussage wird. (In diesem konkreten Fall könnte man $x = -5$ oder $x = 5$ einsetzen.) Je nach Typ der Gleichung gibt es unterschiedliche Algorithmen zum Ermitteln so einer Lösung.

Gleichungssystem Ein *Gleichungssystem* besteht aus mehreren Gleichungen, von denen man möchte, dass sie *allen gemeinsam* wahr werden; z. B. so:

$$\begin{aligned} x^2 - 2 &= 23 \\ y^2 + 10 &= 26 \end{aligned} \tag{22.1}$$

Hier hat man es mit zwei Gleichungen und auch mit zwei Unbekannten zu tun. Das ist aber so noch nicht besonders interessant, weil man in diesem Fall beide Gleichungen separat lösen kann. Interessant wird es, wenn die Gleichungen in irgendeiner Form *gekoppelt* sind, wenn also mindestens eine Unbekannte in mehreren Gleichungen auftaucht:

$$\begin{aligned} x^2 - 2 &= 23 \\ (x - 2) \cdot (x + 3) &= 0 \end{aligned}$$

Hier kommt nur eine Unbekannte vor, aber in beiden Gleichungen. Jede der Gleichungen für sich ist lösbar, aber es gibt keinen Wert, den man für x einsetzen kann, der beide Gleichungen simultan erfüllt. (Warum?) Man würde also hier sagen, dass das System keine Lösung hat. Noch komplizierter wird es, wenn man mehrere Gleichungen hat, in denen mehrere Unbekannte in mehr als einer Gleichung auftreten.

$$\begin{aligned} x^2 + \sin y - 2^z &= 42 \\ \cos^2 x + \ln z &= 23 \\ 2y^3 + \arctan z^2 &= 0 \end{aligned}$$

Hier ist mit den uns bekannten Mitteln überhaupt nicht ersichtlich, ob es Lösungen gibt und wie man diese gegebenenfalls berechnen könnte.

Aufgabe 22.1. Wie viele verschiedene Lösungen hat das System (22.1)?

Wie angekündigt werden wir uns daher nur mit den einfachsten Gleichungssystemen beschäftigen, die es überhaupt gibt, nämlich den *linearen Gleichungssystemen*, für die wir abkürzend auch *LGS* schreiben werden. Damit ist gemeint, dass solche Systeme nur aus Gleichungen bestehen, die so aussehen:²

LGS

$$3x + 5y - 4z = 23$$

Auf einer Seite der Gleichung tauchen keine Variablen auf. Auf der anderen Seite steht eine Summe von Termen und in jedem dieser Terme kommt genau *eine Variable* vor, und zwar in der Form eines Produktes aus dieser Variablen und einem konstanten Faktor.³ Es sind *keinerlei* kompliziertere Ausdrücke zugelassen; also weder höhere Potenzen (x^3) noch trigonometrische Funktionen ($\sin y$) noch Produkte mehrerer Variablen (xy) noch andere „wilde“ Konstruktionen. Zugelassen sind allerdings beliebig viele Variablen und beliebig viele Gleichungen. Hier ein Beispiel mit zwei Gleichungen und drei Unbekannten, an dem man auch gleich sieht, dass die Anzahl der Gleichungen und die Anzahl der Unbekannten nicht übereinstimmen muss.

$$3x + 5y - 4z = 23$$

$$1/3 \cdot z - \pi x = 42$$

Die Faktoren vor den Variablen nennt man ihre *Koeffizienten*.

Koeffizient

Für eine systematische Vorgehensweise ist es sehr hilfreich, so ein System derart zu vereinheitlichen, dass immer *alle* Variablen in allen Gleichungen in derselben Reihenfolge vorkommen:

$$3 \cdot x + 5 \cdot y + (-4) \cdot z = 23$$

$$-\pi \cdot x + 0 \cdot y + 1/3 \cdot z = 42$$

Die Koeffizienten sind hier grün hervorgehoben. Wir haben die Variable y in die zweite Gleichung „hineingeschummelt“, indem wir ihr den Koeffizienten null spendiert haben.

Da die Namen der Variablen keine wesentliche Rolle spielen, braucht man zur Darstellung eines linearen Gleichungssystems eigentlich nur die Koeffizienten und die rechten Seiten der Gleichungen. Man stellt lineare Gleichungssysteme daher abkürzend durch *Matrizen* dar, um sich Schreibaarbeit zu sparen. Für das obige System sähe das so aus:⁴

$$\left(\begin{array}{ccc|c} 3 & 5 & -4 & 23 \\ -\pi & 0 & 1/3 & 42 \end{array} \right)$$

Man nennt das die *erweiterte Koeffizientenmatrix*. Schreibt man nur die Koeffizienten auf, spricht man einfach von der *Koeffizientenmatrix*.

Koeffizientenmatrix

$$\left(\begin{array}{ccc} 3 & 5 & -4 \\ -\pi & 0 & 1/3 \end{array} \right)$$

²Oder die man zumindest so umformen kann, dass sie derart aussehen.

³Trotz des Minuszeichens ist das eine Summe. Der letzte Faktor ist die negative Zahl -4 .

⁴Die senkrechte Linie ist nur eine Hilfslinie und nicht Bestandteil der Matrix.

homogen Ein lineares Gleichungssystem, bei dem die rechte Spalte der erweiterten Koeffizientenmatrix nur aus Nullen besteht, nennt man *homogen*.

Aufgabe 22.2. Geben Sie die erweiterte Koeffizientenmatrix dieses Systems an.

$$\begin{aligned} 6x - y - z &= 4 \\ x + y + 10z &= -6 \\ 2x - y + z &= -2 \end{aligned}$$

Aufgabe 22.3. Sei $\mathbf{M}$ die folgende Matrix und $\mathbf{x}$ der Vektor (x, y, z) . Berechnen Sie das Produkt $\mathbf{M} \cdot \mathbf{x}$.

$$\mathbf{M} = \begin{pmatrix} 6 & -1 & -1 \\ 1 & 1 & 10 \\ 2 & -1 & 1 \end{pmatrix}$$

Das Standardlösungsverfahren für lineare Gleichungssysteme ist das sogenannte *Gaußsche Eliminationsverfahren*. Es beruht darauf, dass manche lineare Gleichungssysteme besonders einfach zu lösen sind. Zum Beispiel dieses hier:

$$\begin{aligned} 2x + 5y - 2z &= -3 \\ 3y + 2z &= 7 \\ 5z &= 10 \end{aligned} \tag{22.2}$$

Hier kann man an der letzten Gleichung erkennen, dass z den Wert 2 haben muss. Setzt man das in die zweite Gleichung ein, vereinfacht sich diese zu $3y + 4 = 7$ bzw. $3y = 3$, woran man sofort sieht, dass für y als Wert nur 1 infrage kommt. Setzt man die beiden Werte dann in die erste Gleichung ein, so erhält man nach einer einfachen Umformung $x = -2$. Die einzige mögliche Lösung ist also: $x = -2$, $y = 1$, $z = 2$.

Lösung An dieser Stelle halten wir gleich mal fest, dass wir unter einer *Lösung* eines Gleichungssystems mit n Unbekannten immer ein *n-Tupel* verstehen, aus dem für alle Variablen hervorgeht, welche Variable welchen Wert bekommt.⁵ (Siehe dazu auch Aufgabe 22.3.) Die einzige Lösung in diesem Fall wäre also das Tupel $(-2, 1, 2)$. Gleichungssysteme können aber auch mehrere Lösungen, also eine *Menge* von Lösungen haben; siehe Aufgabe 22.1. Insbesondere kann die Lösungsmenge auch die leere Menge sein, wenn es gar keine Lösung gibt.

Zurück zum System (22.2). Warum ging das so einfach? Weil die Gleichungen schön brav nach dem folgendem Muster arrangiert waren: Die Gleichungen weiter unten enthielten immer *weniger* Variablen als die über ihnen. Besonders deutlich wird das an der erweiterten Koeffizientenmatrix:

$$\left(\begin{array}{ccc|c} \textcircled{2} & 5 & -2 & -3 \\ 0 & \textcircled{3} & 2 & 7 \\ 0 & 0 & \textcircled{5} & 10 \end{array} \right)$$

⁵Nachdem man sich irgendwie darauf geeinigt hat, in welcher Reihenfolge die Variablen aufgeschrieben werden. In unserem Fall wird die Reihenfolge implizit dadurch festgelegt, wie wir die Gleichungen aufgeschrieben haben.

Den ersten Eintrag einer Zeile, der *nicht* null ist, nennt man auch ihren *Leitkoeffizienten*. Die Leitkoeffizienten wurden oben durch Kreise markiert. Steht in einer Matrix jeder Leitkoeffizient weiter rechts als der über ihm und stehen alle Zeilen ohne Leitkoeffizient (also solche, in denen nur Nullen stehen) ganz unten, so sagt man, dass die Matrix *Zeilenstufenform* (engl. *row echelon form*) hat, abgekürzt *ZSF*. Die Leitkoeffizienten müssen dafür nicht zwangsläufig alle auf der Hauptdiagonalen stehen. Diese Matrix hat z. B. auch Zeilenstufenform:

Leitkoeffizient

Zeilenstufenform

$$\begin{pmatrix} \textcircled{2} & 5 & 8 & -2 \\ 0 & 0 & \textcircled{3} & 2 \\ 0 & 0 & 0 & \textcircled{5} \end{pmatrix}$$

Aufgabe 22.4. Geben Sie die Lösungsmenge des folgenden Gleichungssystems an:

$$\begin{aligned} -x + 3y + z &= 13 \\ 2y + z &= 7 \\ 7z &= -21 \end{aligned}$$

Aufgabe 22.5. Eine quadratische Matrix heißt *obere Dreiecksmatrix*, wenn alle Einträge unterhalb der Hauptdiagonalen verschwinden. Geben Sie eine 3×3 -Matrix an, die obere Dreiecksmatrix ist, aber nicht Zeilenstufenform hat.

Dreiecksmatrix

Die Idee des Gaußverfahrens besteht darin, lineare Gleichungssysteme so umzuformen, dass sich einerseits eine Zeilenstufenform ergibt, sich dabei andererseits aber nicht die Lösungsmenge ändert. (Denn was hat man davon, wenn man ein leicht zu lösendes Gleichungssystem konstruiert, dessen Lösung aber mit dem Ausgangssystem nichts mehr zu tun hat?)

Dabei geht man Schritt für Schritt vor und wendet in jedem Schritt eine sogenannte *elementare Zeilenumformung* (EZ) auf die erweiterte Koeffizientenmatrix an. Davon gibt es drei Typen:

elementare
Zeilenumformung

- (I) Man darf eine Zeile der Matrix mit einem Faktor multiplizieren. Das entspricht der Multiplikation von beiden Seiten einer Gleichung mit demselben Faktor. Sie wissen aus der Schule, dass man das mit einer Gleichung machen darf, ohne die Aussage der Gleichung zu ändern. Multipliziert man z. B. $3x - 4y = 1$ mit 2, so erhält man $6x - 8y = 2$.
Allerdings darf der Faktor natürlich *nicht* die Null sein, denn sonst macht man aus einer Gleichung wie $3x - 4y = 1$, die etwas darüber aussagt, wie x und y zusammenhängen, die Gleichung $0 = 0$, die gar nichts aussagt.
- (II) Man darf zwei Zeilen vertauschen. Die Lösungsmenge eines Systems von Gleichungen hängt sicher nicht davon ab, in welcher Reihenfolge die Gleichungen aufgeschrieben werden.
- (III) Man darf das Vielfache einer Zeile zu einer anderen addieren. Bei diesem Typ von elementarer Zeilenumformung ist vielleicht am wenigsten klar, warum sich die Lösungsmenge nicht ändert. Ein Beispiel:

$$3x + 2y = 8 \tag{1}$$

$$7x - 4y = 5 \quad (2)$$

Wenn beide Gleichungen erfüllt sind, ist sicher auch diese Gleichung erfüllt, die sich durch Multiplikation der ersten mit dem Faktor 2 ergibt:

$$6x + 4y = 16 \quad 2 \cdot (1)$$

Wenn aber die beiden Gleichungen (2) und $2 \cdot (1)$ erfüllt sind, wenn also jeweils links und rechts vom Gleichheitszeichen dasselbe steht, dann gilt das auch für die Summe der beiden:

$$13x = 21 \quad 2 \cdot (1) + (2)$$

Eine elementare Zeilenumformung vom Typ (III) ersetzt also (z. B.) die Gleichungen (1) und (2) durch diese beiden:⁶

$$3x + 2y = 8 \quad (1)$$

$$13x = 21 \quad (2')$$

Und wir haben uns überlegt, dass diese beiden erfüllt sind, wenn die beiden Ausgangsgleichungen erfüllt sind. Aber gilt das auch umgekehrt? Oder haben wir durch diese Umformung evtl. neue Lösungen „dazugeschummelt“? Wir müssen uns noch überlegen, dass jede Lösung der beiden Gleichungen (1) und (2') auch eine Lösung von (1) und (2) ist. Das liegt daran, dass man aus (1) und (2') die ursprünglichen Gleichungen mit derselben Methode wieder rekonstruieren kann: Ersetze (2') durch die Gleichung $(2') - 2 \cdot (1)$. Probieren Sie's aus!

Mit der richtigen Strategie können wir nun durch sukzessive Anwendung von elementaren Zeilenumformungen jede Matrix in Zeilenstufenform umformen. Das führen wir an einem Beispiel vor.

$$\begin{pmatrix} 0 & 0 & -1 & 5 \\ 2 & -2 & 1 & 3 \\ 4 & -4 & 3 & 8 \end{pmatrix}$$

Man arbeitet immer *von oben nach unten* und gleichzeitig *von links nach rechts*. In der Zeile, die gerade „dran“ ist, bringt man zunächst einen Leitkoeffizienten an eine Stelle, die möglichst weit links ist.

Wir fangen also mit der ersten Zeile an. Der Leitkoeffizient wäre die Zahl -1 . Wir können aber durch Vertauschen mit der zweiten Zeile erreichen, dass wir einen Leitkoeffizienten weiter links haben.

$$\begin{pmatrix} \textcircled{2} & -2 & 1 & 3 \\ 0 & 0 & -1 & 5 \\ 4 & -4 & 3 & 8 \end{pmatrix}$$

⁶Durch diesen Typ Umformung ändert sich also nicht die *Anzahl* der Gleichungen.

Den so erhaltenen Koeffizienten (oben eingekreist) nennt man auch das *Pivotelement*.⁷ Nun sorgt man dafür, dass unter dem Pivotelement nur Nullen stehen. Eine Null (in der zweiten Zeile) haben wir schon, aber in der dritten Zeile steht an der entsprechenden Stelle eine Vier. Solche „störenden“ Werte entfernt man dadurch, dass man ein passendes Vielfaches der Zeile mit dem Pivotelement addiert. Hier multiplizieren wir die erste Zeile mit dem Faktor -2 und addieren das Ergebnis zur dritten.

Pivotelement

$$\begin{pmatrix} 2 & -2 & 1 & 3 \\ 0 & 0 & \textcircled{-1} & 5 \\ 0 & 0 & 1 & 2 \end{pmatrix}$$

Wir können uns jetzt der nächsten Zeile zuwenden. Das Schöne am Gaußverfahren ist dabei, dass Zeilen und Spalten, die schon fertig bearbeitet wurden, im Rest des Verfahrens keine Rolle mehr spielen. Das wird oben durch den grauen Hintergrund angedeutet. Diese Komponenten der Matrix sind quasi „inaktiv“: sie ändern sich nicht mehr und werden auch nicht mehr für elementare Zeilenumformungen verwendet. Der Teil der Matrix, den man noch modifizieren muss, wird also sukzessive immer kleiner.

Diesmal müssen wir keine Zeilen vertauschen, um ein Pivotelement für die nächste Zeile zu finden. Wir haben es oben schon eingekreist. Um unter dieser Komponente Nullen zu generieren, addieren wir einfach die zweite Zeile zur dritten.⁸

$$\begin{pmatrix} 2 & -2 & 1 & 3 \\ 0 & 0 & -1 & 5 \\ 0 & 0 & 0 & 7 \end{pmatrix}$$

Wir haben nun eine Zeilenstufenform erreicht.

Aufgabe 22.6. Bringen Sie die Matrix aus Aufgabe 22.2 auf Zeilenstufenform und lösen Sie damit dann das zugehörige lineare Gleichungssystem.

Aufgabe 22.7. Bringen Sie die folgende Matrix auf Zeilenstufenform:

$$\left(\begin{array}{ccc|c} 1 & 0 & 1 & 5 \\ 2 & 1 & 3 & 8 \\ 3 & 1 & 4 & 12 \end{array} \right)$$

Aufgabe 22.8. Bringen Sie die folgende Matrix auf Zeilenstufenform:

$$\left(\begin{array}{ccc|c} 3 & 1 & 5 & 0 \\ 2 & 1 & -2 & 6 \\ -1 & -1 & 9 & -12 \end{array} \right)$$

⁷Vom französischen Wort *pivot* für Dreh- oder Angelpunkt. Falls Sie das Gaußverfahren in ein Computerprogramm übersetzen, ist es übrigens nicht völlig egal, *welches* Pivotelement Sie wählen. (Wir hätten hier ja auch die Vier aus der dritten Zeile nehmen können.) Man möchte nämlich die unvermeidlichen **Rundungsfehler** möglichst minimieren.

⁸Beachten Sie, dass sich links vom Pivotelement durch diese elementare Zeilenumformung nichts mehr ändert, weil dort nur Nullen stehen.

Wir können mit dieser Methode jede Matrix auf Zeilenstufenform bringen. Das heißt aber nicht, dass wir jedes lineare Gleichungssystem lösen können. Die Matrix, die sich bei Aufgabe 22.7 ergeben hat, sieht z. B. so aus:

$$\left(\begin{array}{ccc|c} 1 & 0 & 1 & 5 \\ 0 & 1 & 1 & -2 \\ 0 & 0 & 0 & -1 \end{array} \right)$$

Wenn Sie die letzte Zeile wieder in eine Gleichung zurückübersetzen, dann lautet diese wie folgt:

$$0x + 0y + 0z = -1 \quad (22.3)$$

Es ist aber offensichtlich, dass auf der linken Seite dieser Gleichung immer null steht. Welche Werte Sie für x , y und z auch immer einsetzen, es kann niemals -1 herauskommen. Daher ist dieses Gleichungssystem *nicht* lösbar: die Lösungsmenge ist $\emptyset$.

Es kann aber auch noch etwas anderes passieren. Dafür sehen wir uns die Matrix an, die sich in Aufgabe 22.8 ergeben hat:

$$\left(\begin{array}{ccc|c} -1 & -1 & 9 & -12 \\ 0 & -1 & 16 & -18 \\ 0 & 0 & 0 & 0 \end{array} \right)$$

Die letzte Zeile sieht als Gleichung so aus:

$$0x + 0y + 0z = 0$$

Im Gegensatz zu Gleichung (22.3) ist diese Gleichung nicht unlösbar. Es ist vielmehr so, dass sie *immer* lösbar ist, unabhängig von den eingesetzten Werten. Mit anderen Worten: sie sagt über die Variablen nichts aus.

Daher können wir keinen bestimmten Wert für z in die zweite Gleichung einsetzen, sondern erhalten einfach dies:

$$-y + 16z = -18$$

Lösen wir nach y auf, so ergibt sich $y = 18 + 16z$. Wir können das so interpretieren, dass y von z abhängt: wir können z frei wählen, und daraus ergibt sich zwangsläufig ein Wert für y . Setzen wir das in die erste Gleichung ein:

$$-x - (18 + 16z) + 9z = -12$$

Auch das können wir auflösen: $x = -6 - 7z$. x und y hängen also beide von der Wahl des Wertes für z ab und z kann irgendeine reelle Zahl sein. Die Lösungsmenge sieht daher so aus:

$$\{(-6 - 7z, 18 + 16z, z) : z \in \mathbb{R}\} \quad (22.4)$$

Wir haben unendlich viele Lösungen.

Aufgabe 22.9. Erinnert Sie die Menge in (22.4) an etwas?

★**Aufgabe 22.10.** In Aufgabe 22.6 wurden elementare Zeilenumformungen vom Typ (I) verwendet, um Brüche zu vermeiden. Es ist aber wohl offensichtlich, dass man im Prinzip auf solche elementaren Zeilenumformungen völlig verzichten könnte. Elementare Zeilenumformungen vom Typ (II) haben wir dann benötigt, wenn wir Pivotelemente an die richtige Stelle bringen wollten. Kann man zur Not auf die auch verzichten und nur mit elementaren Zeilenumformungen vom Typ (III) auskommen?

Das bisher Erreichte können wir wie folgt zusammenfassen:

- Wenn wir die erweiterte Koeffizientenmatrix eines linearen Gleichungssystems auf Zeilenstufenform gebracht haben, können wir zunächst alle Zeilen ignorieren, in denen ausschließlich Nullen stehen.
- Ist die letzte der verbleibenden Zeilen eine, in der links vom Gleichheitszeichen nur Nullen stehen (und rechts logischerweise dann nicht), so ist das Gleichungssystem nicht lösbar.
- Tritt der vorherige Fall nicht ein und gibt es mehr Unbekannte als Zeilen,⁹ so können wir alle Variablen, in deren Spalte kein Leitkoeffizient steht, frei wählen. Wir erhalten dann durch Rückwärtseinsetzen eine unendliche Menge von Lösungen, die von ein oder mehreren Parametern abhängen.
- Anderenfalls ergibt sich durch Rückwärtseinsetzen eine eindeutig bestimmte Lösung.

Aufgabe 22.11. Geben Sie die Lösungsmenge des folgenden LGS an.

$$\begin{aligned}x_1 + x_2 + 4x_3 + 2x_4 &= 7 \\ -x_1 - 4x_2 + 4x_3 + 3x_4 &= 13 \\ -x_2 - 4x_3 - 5x_4 &= 0 \\ -5x_1 - x_2 - 3x_3 + 3x_4 &= -26\end{aligned}$$

Aufgabe 22.12. Geben Sie die Lösungsmenge des folgenden LGS in $\mathbb{Z}_7$ (!) an.

$$\begin{aligned}3x_1 + 4x_2 + 3x_3 &= 2 \\ 4x_1 + 2x_2 + 2x_3 &= 2 \\ 6x_1 + x_2 + 2x_3 &= 6\end{aligned}$$

Gehen Sie dabei genauso vor wie bisher. Der Unterschied ist lediglich, dass alle Zahlen und Variablen als Elemente von $\mathbb{Z}_7$ betrachtet werden.

Aufgabe 22.13. Ermitteln Sie die eindeutige Lösung des folgenden Gleichungssystems, wobei a eine reelle Konstante sein soll.

$$\begin{aligned}ax + y + z &= 1 + a \\ x + z &= 2 \\ x - z &= 0\end{aligned}$$

⁹Man nennt das Gleichungssystem in so einem Fall *unterbestimmt*.

Nachdem man ein lineares Gleichungssystem auf Zeilenstufenform gebracht hat, kann man es optional noch weiter mithilfe von elementaren Zeilenumformungen modifizieren. Das Ziel dabei ist, ein Gleichungssystem zu erhalten, aus dem die Lösung direkt – ohne Rückwärtseinsetzen – ablesbar ist. Man spricht dann vom *Gauß-Jordan-Algorithmus*, der nach dem deutschen Geodäten [Wilhelm Jordan](#) benannt ist, der das Verfahren jedoch nicht entwickelt, sondern nur popularisiert hat. Dabei geht man so vor:

- Wir haben bereits ein System in Zeilenstufenform. Unser Ziel ist es, in den Spalten oberhalb der Leitkoeffizienten Nullen zu erzeugen.
- Während wir beim Erzeugen der Zeilenstufenform von oben nach unten und von links nach rechts vorgegangen sind, arbeiten wir uns nun von rechts nach links und von unten nach oben durch.
- Es werden keine Zeilen mehr vertauscht. Der Leitkoeffizient, der gerade „dran“ ist, ist zwangsläufig das Pivotelement.
- All diese Pivotelemente werden als Erstes (durch Multiplikation der Zeile mit dem Kehrwert) auf den Wert eins gebracht.

Schauen wir uns als Beispiel die Matrix in Zeilenstufenform an, die wir in Aufgabe 22.6 errechnet haben. Der Leitkoeffizient ganz unten ist der, mit dem wir beginnen.

$$\left(\begin{array}{ccc|c} 1 & 1 & 10 & -6 \\ 0 & -7 & -61 & 40 \\ 0 & 0 & \textcircled{50} & -50 \end{array} \right)$$

Wir dividieren durch 50:

$$\left(\begin{array}{ccc|c} 1 & 1 & 10 & -6 \\ 0 & -7 & -61 & 40 \\ 0 & 0 & \textcircled{1} & -1 \end{array} \right)$$

Nun wollen wir aus den Zahlen oberhalb des Pivotelements Nullen machen. Dazu subtrahieren wir das Zehnfache der dritten Zeile von der ersten und addieren das 61fache zur zweiten. Das nächste Pivotelement ist der Leitkoeffizient der zweitletzten Zeile.

$$\left(\begin{array}{ccc|c} 1 & 1 & 0 & 4 \\ 0 & \textcircled{-7} & 0 & -21 \\ 0 & 0 & 1 & -1 \end{array} \right)$$

Beachten Sie, dass sich links vom Pivotelement nichts mehr ändert, weil die Matrix ja schon in Zeilenstufenform vorlag und unser Pivotelement ein Leitkoeffizient war. Wir dividieren nun durch -7 :

$$\left(\begin{array}{ccc|c} 1 & 1 & 0 & 4 \\ 0 & \textcircled{1} & 0 & 3 \\ 0 & 0 & 1 & -1 \end{array} \right)$$

Um aus der Zahl oberhalb des aktuellen Pivotelements eine Null zu machen, subtrahieren wir die zweite von der ersten Zeile.

$$\left(\begin{array}{ccc|c} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 3 \\ 0 & 0 & 1 & -1 \end{array} \right) \quad (22.5)$$

Auch hier sorgt das Verfahren automatisch dafür, dass wir immer nur in der Spalte etwas ändern, in der sich unser Pivotelement befindet, sowie ggf. in der Spalte ganz rechts. Schreiben wir (22.5) wieder als Gleichungssystem:

$$\begin{aligned} x &= 1 \\ y &= 3 \\ z &= -1 \end{aligned}$$

Damit sollte klar sein, was das Ziel des Gauß-Jordan-Verfahrens ist: am Ende steht die Lösung einfach da und man hat keine weitere Arbeit.

Natürlich muss man zusätzliche Arbeit investieren, um so weit zu kommen. Wenn man also selbst mit Bleistift und Papier rechnet und nur *ein* lineares Gleichungssystem lösen will, ist es wahrscheinlich einfacher, nur bis zur Zeilenstufenform umzuformen und dann durch Rückwärtseinsetzen zu lösen. Das Gauß-Jordan-Verfahren ist erst dann interessant, wenn man es mit *mehreren* Systemen zu tun hat, die sich nur durch ihre rechten Seiten unterscheiden. Man führt das Verfahren dann gleichzeitig mit allen rechten Seiten durch und löst alle Gleichungssysteme auf einmal. Außerdem ist dieses Verfahren auch in anderen Situationen nützlich. Wir werden z. B. sehen, dass man damit **Inverse von Matrizen** ausrechnen kann.

Zum Schluss überlegen wir uns noch, dass sich alle elementaren Zeilenumformungen als Matrizenmultiplikationen darstellen lassen. Das werden wir nämlich demnächst brauchen. Beginnen wir mit Zeilenumformungen vom Typ (III). Um z. B. zur zweiten Zeile einer $4 \times n$ -Matrix **A** das Dreifache der vierten Zeile von **A** zu addieren, kann man das hier berechnen:

$$\begin{array}{c} \text{2. Zeile} \rightarrow \left(\begin{array}{cccc} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 3 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{array} \right) \cdot \mathbf{A} \\ \qquad \qquad \qquad \downarrow \text{4. Spalte} \end{array}$$

Und so kann man jede elementare Zeilenumformung vom Typ (III) durch Multiplikation mit einer entsprechenden Matrix von links erreichen. Probieren Sie's selbst anhand von Beispielen aus!

★**Aufgabe 22.14.** Überlegen Sie sich, dass man auch elementare Zeilenumformungen der anderen beiden Typen durch Multiplikation mit einer Matrix von links realisieren kann. (Hinweis: In beiden Fällen verwendet man dazu wie oben Matrizen, die der Einheitsmatrix sehr ähnlich sind.)

Bemerkungen

Prinzipiell können Sie natürlich lineare Gleichungssysteme lösen, wie Sie wollen, und manchmal bieten sich andere Verfahren oder Abkürzungen an. Das hier vorgestellte Gaußverfahren hat allerdings den Vorteil, für beliebige Matrizen das effizienteste zu sein. In den folgenden Abschnitten werden wir außerdem sehen, dass man es auch für andere Zwecke einsetzen kann.

★ Lineare Gleichungssysteme im Computer

Stellt man Matrizen so dar wie [im letzten Abschnitt](#), so kann man Zeilenumformungen vom Typ (III) wie folgt implementieren:

```
def rowMod (M,i,j,x):
    M[i] = [a + x * b for a, b in zip(M[i], M[j])]
```

Die Funktion addiert zur i -ten Zeile der Matrix M das x -fache der j -ten Zeile. Man beachte allerdings, dass `rowMod` die Matrix modifiziert!

Die Umwandlung einer Matrix in Zeilenstufenform lässt sich damit z. B. so implementieren:

```
def rowEchelon (M):
    row, col = 0, 0
    rows, cols = len(M), len(M[0])
    while row < rows and col < cols:
        # try to find pivot element
        if M[row][col] == 0:
            for r in range(row + 1, rows):
                if M[r][col] != 0:
                    rowMod(M, row, r, 1)
                    break
        # no pivot element, skip column
        if M[row][col] == 0:
            col += 1
            continue
        # pivot element found, create zeros
        pivot = M[row][col]
        for r in range(row + 1, rows):
            if M[r][col] != 0:
                rowMod(M, r, row, -M[r][col] / pivot)
        # next row, next column
        row += 1
        col += 1
```


Aber Achtung! Das ist nur ein *Proof of Concept* und kein Programm, das man in der Praxis verwenden sollte. Wenn Sie beispielsweise

```
M = [[6, -1, -1, 4], [1, 1, 10, -6], [2, -1, 1, -2]]
rowEchelon(M)
```

eingeben, werden Sie sehen, dass M danach *nicht* in Zeilenstufenform ist. Warum ist das so? Siehe dazu den [Abschnitt über Fehlerfortpflanzung im Computer](#).

23. Lineare Abbildungen

Nachdem wir nun das nötige Handwerkszeug zusammen haben, wenden wir uns wieder der Geometrie zu. Eine Funktion f von $\mathbb{R}^2$ nach $\mathbb{R}^2$ bildet ganz allgemein Punkte auf Punkte ab. Damit bildet sie auch geometrische Figuren auf geometrische Figuren ab: Ist D beispielsweise die Menge der Punkte eines Dreiecks, so ist $f[D] = \{f(\mathbf{p}) : \mathbf{p} \in D\}$, das sogenannte *Bild* von D unter f ebenfalls eine Punktmenge. Wenn f irgendeine Funktion ist, kann $f[D]$ eine beliebig komplizierte Punktmenge sein, die wir gar nicht mehr als „Figur“ identifizieren würden. Wir wollen uns im Folgenden auf möglichst einfache Abbildungen beschränken, die sich geometrisch „sinnvoll“ verhalten. Insbesondere fordern wir zwei Eigenschaften von so einer Funktion f , damit wir sie als einfach betrachten:

Bild $f[D]$

- Ist $\mathbf{v}$ der Vektor, der den Punkt $\mathbf{p}$ auf den Punkt $\mathbf{q}$ verschiebt, dann soll $f(\mathbf{v})$ der Vektor sein, der den Punkt $f(\mathbf{p})$ auf den Punkt $f(\mathbf{q})$ verschiebt. Mit anderen Worten: Wir möchten f interpretieren können als Funktion, die Vektoren auf Vektoren abbildet und dabei die Beziehungen zwischen Ortsvektoren und verschiebenden Vektoren erhält.
- Ist $\mathbf{w}$ das Produkt von $\mathbf{v}$ mit dem Skalar α , dann soll $f(\mathbf{w})$ das Produkt von $f(\mathbf{v})$ mit α sein. Mit anderen Worten: f soll Parallelität von Vektoren erhalten und auch die Längenverhältnisse von parallelen Vektoren zueinander.

Wir fassen das in einer Definition ganz allgemein zusammen.

Eine Funktion $f: \mathbb{R}^m \rightarrow \mathbb{R}^n$ wird **linear** genannt, wenn für alle $\mathbf{v}, \mathbf{w} \in \mathbb{R}^n$ und alle $\alpha \in \mathbb{R}$ die beiden folgenden Beziehungen gelten:

$$(L1) \quad f(\mathbf{v} + \mathbf{w}) = f(\mathbf{v}) + f(\mathbf{w})$$

$$(L2) \quad f(\alpha \cdot \mathbf{v}) = \alpha \cdot f(\mathbf{v})$$

Traditionell spricht man eher von linearen *Abbildungen* als von Funktionen.

Definition

Sei nun f irgendeine lineare Abbildung und seien $\mathbf{p}$ und $\mathbf{q}$ zwei verschiedene Punkte aus dem Definitionsbereich von f . Erstens haben wir dann

$$f(\mathbf{0}) = f(0 \cdot \mathbf{p}) = 0 \cdot f(\mathbf{p}) = \mathbf{0}$$

wegen (L2). Zweitens gibt – siehe (20.4) – für jeden Punkt $\mathbf{x}$ auf der Geraden durch $\mathbf{p}$ und $\mathbf{q}$ Skalare α und β mit $\mathbf{x} = \alpha\mathbf{p} + \beta\mathbf{q}$ sowie $\alpha + \beta = 1$. Mit (L1) und (L2) folgt $f(\mathbf{x}) = \alpha f(\mathbf{p}) + \beta f(\mathbf{q})$. $f(\mathbf{x})$ liegt also auf der Geraden durch $f(\mathbf{p})$ und $f(\mathbf{q})$. Und weil

sich α und β nicht geändert haben, liegt $f(\mathbf{x})$ dort im selben Verhältnis zu diesen beiden Punkten wie $\mathbf{x}$ zu $\mathbf{p}$ und $\mathbf{q}$.

Aufgabe 23.1. In der Argumentation des letzten Absatzes fehlt ein Spezialfall. Sehen Sie, was vergessen wurde?

Insgesamt können wir zusammenfassen:

Lemma 23.1

Lineare Abbildungen bilden Nullvektoren auf Nullvektoren ab und Geraden auf Geraden oder Punkte. Wird eine Gerade auf eine Gerade abgebildet, so bleiben die Streckenverhältnisse auf dieser Geraden erhalten. Werden zwei parallele Geraden auf zwei Geraden abgebildet, so sind diese Bildgeraden ebenfalls parallel.

Zwei ganz simple Beispiele für lineare Abbildungen sind

$$: \begin{cases} \mathbb{R}^m \rightarrow \mathbb{R}^n \\ \mathbf{v} \mapsto \mathbf{0} \end{cases}$$

und $\text{id}_{\mathbb{R}^m}$. Es ist hoffentlich klar, dass beide die Bedingungen aus der Definition erfüllen (und dass die erste alle Geraden auf einen Punkt „zusammenschnurren“ lässt). Ein weniger triviales Beispiel ist

$$f: \begin{cases} \mathbb{R}^2 \rightarrow \mathbb{R}^2 \\ \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} 2x + 3y \\ 3x - y \end{pmatrix} \end{cases} \quad (23.1)$$

Für f rechnen wir die Bedingung (L1) konkret nach:

$$\begin{aligned} f\left(\begin{pmatrix} v_1 \\ v_2 \end{pmatrix} + \begin{pmatrix} w_1 \\ w_2 \end{pmatrix}\right) &= f\left(\begin{pmatrix} v_1 + w_1 \\ v_2 + w_2 \end{pmatrix}\right) = \begin{pmatrix} 2(v_1 + w_1) + 3(v_2 + w_2) \\ 3(v_1 + w_1) - (v_2 + w_2) \end{pmatrix} \\ &= \begin{pmatrix} 2v_1 + 3v_2 \\ 3v_1 - v_2 \end{pmatrix} + \begin{pmatrix} 2w_1 + 3w_2 \\ 3w_1 - w_2 \end{pmatrix} = f\left(\begin{pmatrix} v_1 \\ v_2 \end{pmatrix}\right) + f\left(\begin{pmatrix} w_1 \\ w_2 \end{pmatrix}\right) \end{aligned}$$

Aufgabe 23.2. Machen Sie das auch für die zweite Bedingung (L2).

Aufgabe 23.3. Berechnen Sie für f aus (23.1) die Funktionswerte für die Argumente $(1, 0)$ und $(0, 1)$.

Mithilfe der definierenden Bedingungen können wir nun eine grundlegende Eigenschaft von linearen Abbildungen herausarbeiten. Dafür führen wir zunächst eine Schreibweise ein.

Mit $\mathbf{e}_k$ bezeichnen wir einen Vektor, dessen k -te Komponente 1 ist und dessen andere Komponenten alle verschwinden. Man nennt ihn den k -ten **kanonischen Einheitsvektor**.

Definition

Je nach Kontext kann $\mathbf{e}_k$ ein Element von $\mathbb{R}^3$, $\mathbb{Q}^7$, $\mathbb{Z}_2^2$ oder von einer anderen Menge von Vektoren sein (wenn k nicht zu groß ist). In $\mathbb{R}^3$ ist $\mathbf{e}_1$ beispielsweise der Vektor $(1, 0, 0)$, in $\mathbb{Z}_7^2$ handelt es sich um $(1, 0)$. In $\mathbb{R}^3$ ist $\mathbf{e}_3$ der Vektor $(0, 0, 1)$, in $\mathbb{R}^2$ ergibt die Bezeichnung $\mathbf{e}_2$ keinen Sinn.

Die Vektoren $\mathbf{e}_1$ bis $\mathbf{e}_n$ nennt man auch die (*kanonischen*) *Basisvektoren* von $\mathbb{R}^n$ (bzw. von K^n wenn K ein anderer Körper wie $\mathbb{Q}$ oder $\mathbb{Z}_5$ ist).

Basisvektor

Mithilfe dieser Schreibweise kann man Vektoren nun gewissermaßen „in ihre Komponenten“ zerlegen, zum Beispiel $(3, 6, -2) = 3 \cdot \mathbf{e}_1 + 6 \cdot \mathbf{e}_2 - 2 \cdot \mathbf{e}_3$.

Nun zu der besagten grundlegenden Eigenschaft. Der Übersichtlichkeit halber werden wir diese nur für das Beispiel $m = 3$ und $n = 2$ betrachten. Sie gilt aber ganz allgemein. Sei also $f: \mathbb{R}^3 \rightarrow \mathbb{R}^2$ linear. Ein Vektor aus $\mathbb{R}^3$ hat ganz allgemein die Form (x_1, x_2, x_3) . Setzt man den in f ein, so erhält man:

$$\begin{aligned} f((x_1, x_2, x_3)) &= f(x_1 \cdot \mathbf{e}_1 + x_2 \cdot \mathbf{e}_2 + x_3 \cdot \mathbf{e}_3) = f(x_1 \cdot \mathbf{e}_1) + f(x_2 \cdot \mathbf{e}_2) + f(x_3 \cdot \mathbf{e}_3) \\ &= x_1 \cdot f(\mathbf{e}_1) + x_2 \cdot f(\mathbf{e}_2) + x_3 \cdot f(\mathbf{e}_3) \end{aligned}$$

Kennt man die drei Funktionswerte $f(\mathbf{e}_1)$ bis $f(\mathbf{e}_3)$, so kennt man also schon alle anderen! Und es kommt noch besser. Dafür schreiben wir diesen Zusammenhang explizit für die Funktion f aus (23.1) auf und verwenden dabei die in Aufgabe 23.3 berechneten Werte:

$$\begin{aligned} f((x_1, x_2)) &= x_1 \cdot f(\mathbf{e}_1) + x_2 \cdot f(\mathbf{e}_2) \\ &= x_1 \cdot \begin{pmatrix} 2 \\ 3 \end{pmatrix} + x_2 \cdot \begin{pmatrix} 3 \\ -1 \end{pmatrix} = \begin{pmatrix} 2 & 3 \\ 3 & -1 \end{pmatrix} \cdot \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \end{aligned}$$

Die Funktion lässt sich als Produkt „Matrix mal Vektor“ darstellen, wobei die Spalten der Matrix die Bilder der Vektoren $\mathbf{e}_1$ und $\mathbf{e}_2$ sind. Das gilt ganz allgemein:

Eine Funktion $f: \mathbb{R}^m \rightarrow \mathbb{R}^n$ ist genau dann linear, wenn sie sich durch eine Matrix darstellen lässt; wenn es also eine Matrix $\mathbf{A} \in \mathbb{R}^{n \times m}$ mit $f(\mathbf{x}) = \mathbf{A} \cdot \mathbf{x}$ für alle $\mathbf{x} \in \mathbb{R}^m$ gibt. Diese Matrix ist eindeutig bestimmt. Ihre Spalten entsprechen (in dieser Reihenfolge) den Vektoren $f(\mathbf{e}_1)$ bis $f(\mathbf{e}_m)$.

Lemma 23.2**Aufgabe 23.4.** Die lineare Abbildung

$$f: \begin{cases} \mathbb{R}^3 \rightarrow \mathbb{R}^2 \\ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \mapsto \begin{pmatrix} 5x - 2y \\ 4x - z \end{pmatrix} \end{cases}$$

soll durch eine Matrix als $f(\mathbf{v}) = \mathbf{A} \cdot \mathbf{v}$ dargestellt werden. Geben Sie die Matrix $\mathbf{A}$ an.

Aufgabe 23.5. Die lineare Abbildung f von $\mathbb{R}^2$ nach $\mathbb{R}^2$ bildet $\mathbf{e}_1$ auf $(2, 1)$ und $\mathbf{e}_2$ auf $(2, -4)$ ab. Geben Sie die zugehörige Matrix und den Vektor $f((3, 2))$ an.

Aufgabe 23.6. Die lineare Abbildung f von $\mathbb{R}^2$ nach $\mathbb{R}^2$ bildet $\mathbf{e}_1$ auf $(2, 1)$ und $2\mathbf{e}_2$ auf $(2, -4)$ ab. Geben Sie den Vektor $f((3, 2))$ an.

Aufgabe 23.7. Die lineare Abbildung f von $\mathbb{R}^2$ nach $\mathbb{R}^2$ bildet $(2, 1)$ auf $(3, 2)$ und $(0, 2)$ auf $(1, -5)$ ab. Geben Sie den Vektor $f((3, 1))$ an. (Hinweis: Überlegen Sie zuerst, wie man $(3, 1)$ in der Form $\alpha \cdot (2, 1) + \beta(0, 2)$ darstellen kann. Wenden Sie dann die definierenden Eigenschaften einer linearen Abbildung an.)

Aufgabe 23.8. Gegeben seien die folgenden $\mathbb{R}^3$ -Vektoren:

$$\mathbf{e}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} \quad \mathbf{e}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \quad \mathbf{v}_1 = \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix} \quad \mathbf{v}_2 = \begin{pmatrix} -4 \\ 5 \\ 7 \end{pmatrix} \quad \mathbf{v}_3 = \begin{pmatrix} 2 \\ -1 \\ 1 \end{pmatrix} \quad \mathbf{v}_4 = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} \quad \mathbf{w} = \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix}$$

Die lineare Abbildung $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ bildet $\mathbf{e}_1$ auf $\mathbf{v}_1$ ab, $\mathbf{e}_2$ auf $\mathbf{v}_2$ und $\mathbf{v}_3$ auf $\mathbf{v}_4$. Geben Sie den Funktionswert $f(\mathbf{w})$ an.

Aufgabe 23.9. Die lineare Abbildung f von $\mathbb{Z}_5^2$ nach $\mathbb{Z}_5^3$ bildet $3\mathbf{e}_1$ auf $(2, 4, 3)$ und $2\mathbf{e}_2$ auf $(1, 0, 3)$ ab. Geben Sie den Vektor $f((4, 1))$ an.

Was passiert, wenn man zwei lineare Abbildungen **komponiert**? Sind f und g zwei solche Funktionen, die durch die Matrizen $\mathbf{A}$ bzw. $\mathbf{B}$ definiert sind, d. h. $f(\mathbf{x}) = \mathbf{A} \cdot \mathbf{x}$ und $g(\mathbf{x}) = \mathbf{B} \cdot \mathbf{x}$, dann ist die Frage also, wie $f \circ g$ aussieht.

Damit man überhaupt sinnvoll von $f \circ g$ reden kann, muss zunächst der Definitionsbereich von f zum Wertebereich von g passen. (Wenn die Bilder von g z. B. Vektoren aus $\mathbb{R}^4$ sind, kann man sie nicht in eine Abbildung f einsetzen, die Vektoren aus $\mathbb{R}^3$ erwartet.) Ist diese Voraussetzung aber erfüllt, so ergibt sich die Komposition ganz einfach:¹⁰

$$(f \circ g)(\mathbf{x}) = f(g(\mathbf{x})) = f(\mathbf{B} \cdot \mathbf{x}) = \mathbf{A} \cdot (\mathbf{B} \cdot \mathbf{x}) = (\mathbf{A} \cdot \mathbf{B}) \cdot \mathbf{x}$$

Aus dieser Gleichung kann man mehrere wichtige Schlüsse ziehen:

- $f \circ g$ lässt sich durch $(f \circ g)(\mathbf{x}) = (\mathbf{A} \cdot \mathbf{B}) \cdot \mathbf{x}$ darstellen. Das bedeutet, dass der Komposition von linearen Abbildungen die Multiplikation der zugehörigen Matrizen entspricht.
- Es beantwortet nebenbei auch die Frage, warum die Multiplikation von Matrizen im Vergleich zur Addition so „umständlich“ definiert ist: weil sie für diese Anwendung „maßgeschneidert“ wurde.
- Aus dem **Kapitel über Funktionen** (siehe Aufgabe 16.23) wissen wir bereits, dass die Komposition von Funktionen im Allgemeinen nicht kommutativ ist. Daher kann die Matrizenmultiplikation natürlich auch nicht kommutativ sein. (Das wissen wir ebenfalls schon, siehe Aufgabe 21.5.)

¹⁰Der letzte Schritt folgt dabei aus der Assoziativität der Matrizenmultiplikation, der ersten Eigenschaft in Lemma 21.2.

- Da $\mathbf{A} \cdot \mathbf{B}$ als Produkt von Matrizen wieder eine Matrix ist, ist die Komposition von linearen Abbildungen wieder eine lineare Abbildung.
- Oben wurde schon erwähnt, dass man nur zueinander „passende“ Abbildungen komponieren kann. Das entspricht der Tatsache, dass man auch nur zueinander „passende“ Matrizen multiplizieren kann.

Wir werden noch sehen, dass man in der Computergrafik häufig eine komplexe Bewegung aus mehreren einfachen linearen Abbildungen zusammensetzt. Wir haben gerade gelernt, dass sich das Ergebnis unabhängig von der Komplexität der Gesamtbewegung durch *eine* Matrix darstellen lässt. Das ist für die Effizienz der Programme sehr wichtig, weil wir dann zur Durchführung der Bewegung nur eine einzige Matrizenmultiplikation durchführen müssen.

Aufgabe 23.10. Wir betrachten die folgenden beiden linearen Abbildungen:

$$f: \begin{cases} \mathbb{R}^3 \rightarrow \mathbb{R}^2 \\ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \mapsto \begin{pmatrix} 2x-3y \\ z-x \end{pmatrix} \end{cases} \quad g: \begin{cases} \mathbb{R}^3 \rightarrow \mathbb{R}^3 \\ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \mapsto \begin{pmatrix} x+2z \\ 5y-z \\ x-y+z \end{pmatrix} \end{cases}$$

Geben Sie die zu $f \circ g$ gehörende Matrix an. Was ist mit $g \circ f$?

Aufgabe 23.11. f und g seien lineare Abbildungen von $\mathbb{R}^2$ nach $\mathbb{R}^2$. f bildet $(2,0)$ auf $(2,4)$ und $(0,1)$ auf $(2,2)$ ab. g ist durch

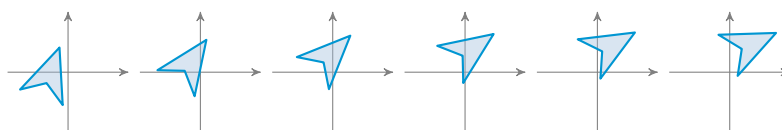
$$g(\mathbf{v}) = \begin{pmatrix} 3 & -1 \\ 2 & 0 \end{pmatrix} \cdot \mathbf{v}$$

definiert. Geben Sie die Matrix zur linearen Abbildung $g \circ f$ an.

24. Geometrische Interpretation linearer Abbildungen

Wir kehren nun wieder zur Geometrie zurück und beschäftigen uns zunächst ganz konkret mit linearen Abbildungen von $\mathbb{R}^2$ nach $\mathbb{R}^2$, die man – wie im letzten Abschnitt angekündigt – interpretieren kann als Funktionen, die Punktmengen auf Punktmengen abbilden.

So werden diese Funktionen auch in der Computergrafik verwendet. Wenn z. B. wie in der folgenden Grafik eine Animation in einem Computer aus vielen Einzelbildern zusammengesetzt ist, dann lässt sich das am effizientesten berechnen, wenn jedes Bild aus dem vorherigen durch eine lineare Abbildung entsteht, weil Grafikprozessoren für solche Aufgaben optimiert sind.

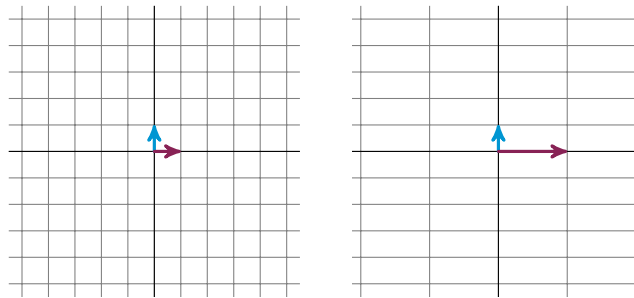


Nach Lemma 23.1 wird das Koordinatengitter, das im Urbildraum aus zwei Scharen äquidistanter paralleler Geraden besteht,¹¹ abgebildet auf zwei andere Scharen paralleler Geraden, die aber nicht mehr unbedingt rechtwinklig zueinander stehen. Außerdem sind die Abstände innerhalb einer der beiden Scharen zwar nach wie vor gleich, sie müssen aber nicht denen in der anderen Schar entsprechen. Anders ausgedrückt werden aus den Quadraten, in die das Koordinatengitter die Ebene zerlegt, kongruente Parallelogramme.

Skalierung Das erste Beispiel, das wir uns anschauen, sind *Skalierungen*. Die folgende Skizze zeigt den Effekt einer Matrix der Form

$$\begin{pmatrix} 2.6 & 0 \\ 0 & 1 \end{pmatrix},$$

die sich nur in einer Komponente von der Einheitsmatrix unterscheidet. Hier wird parallel zur x -Achse um den Faktor 2.6 gestreckt, während die y -Koordinaten sich nicht ändern. Selbstverständlich kann man auch in y -Richtung skalieren oder mit Faktoren aus dem Intervall $(0, 1)$ stauchen.



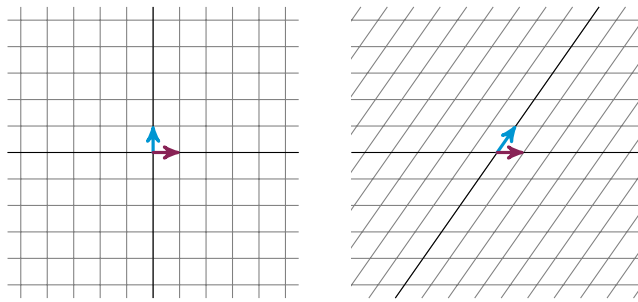
Beachten Sie, dass man anhand der farbig eingezeichneten Bilder der beiden Basisvektoren den ganzen Rest konstruieren kann und dass diese Bilder in den Spalten der Matrix stehen. Beides entspricht dem, was auf den vorherigen Seiten theoretisch erarbeitet wurde, und gilt natürlich auch für alle folgenden Beispiele.

Scherung Das zweite Beispiel sind *Scherungen*. Die Matrix zur Skizze ist

$$\begin{pmatrix} 1 & 0.7 \\ 0 & 1 \end{pmatrix}. \quad (24.1)$$

Auch hier unterscheidet sich nur eine Komponente von der Einheitsmatrix, allerdings wurde in diesem Fall eine der Nullen modifiziert. Die Punkte behalten ebenfalls ihre y -Koordinate, anders als bei der Skalierung hängt allerdings die horizontale Verschiebung von ihrem Abstand von der x -Achse ab – und davon, ob sie oberhalb oder unterhalb derselben liegen.

¹¹Nämlich aus den waagerechten Geraden der Form $(0, z) + \mathbb{R}(1, 0)$ für $z \in \mathbb{Z}$ und den senkrechten der Form $(z, 0) + \mathbb{R}(0, 1)$ für $z \in \mathbb{Z}$.

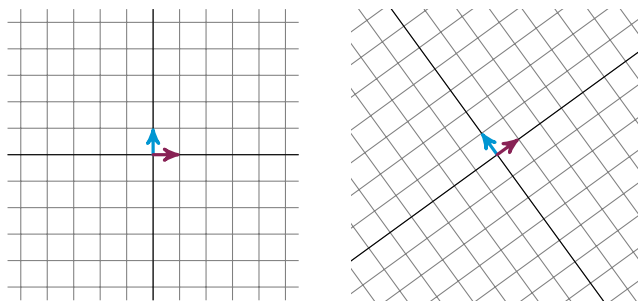


Das dritte Beispiel zeigt eine *Drehung* um den Nullpunkt um den Winkel $\pi/5$. Die Matrix dazu ist

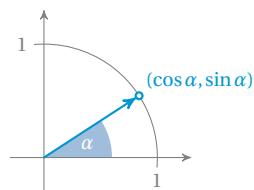
Drehung

$$\begin{pmatrix} \cos(\pi/5) & -\sin(\pi/5) \\ \sin(\pi/5) & \cos(\pi/5) \end{pmatrix}.$$

Eine lineare Abbildung kann offenbar nur Drehungen um den Nullpunkt darstellen, weil dieser sich nach Lemma 23.1 nicht bewegen darf.



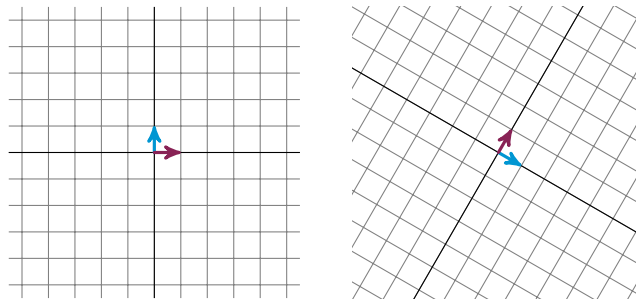
Die erste Spalte der Matrix zeigt den entsprechend gedrehten Basisvektor $\mathbf{e}_1$, den man mithilfe der *trigonometrischen Funktionen* berechnen kann. Die zweite Spalte kann man ebenso aus dem zweiten Basisvektor erhalten oder durch Drehung des ersten Funktionswertes um 90 Grad, wie wir es uns schon bei der *Multiplikation komplexer Zahlen* überlegt hatten.



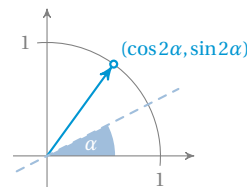
Das vierte Beispiel sieht auf den ersten Blick wie eine Drehung aus, allerdings erkennt man bei genauerem Hinsehen, dass die Bilder der Basisvektoren quasi ihre Rolle getauscht haben. Es handelt sich um eine *Spiegelung* an einer durch den Nullpunkt verlaufenden Achse. Die Matrix in diesem Fall ist

Spiegelung

$$\begin{pmatrix} \cos(\pi/3) & \sin(\pi/3) \\ \sin(\pi/3) & -\cos(\pi/3) \end{pmatrix}.$$



Wieso kommen auch hier Kosinus, Sinus und Winkel vor? Die Achse, an der gespiegelt wird, ist die Gerade durch den Nullpunkt, die man erhält, wenn man die horizontale Achse um $\pi/6$ dreht. Der in der Matrix verwendete Winkel ist das Doppelte davon. Mithilfe einer kleinen Skizze kann man sich leicht davon überzeugen, wieso die erste Spalte der Matrix so aussehen muss. Die zweite ergibt sich dann durch Drehung der ersten um -90° . (Man beachte das Vorzeichen.)



Für das fünfte Beispiel gibt es keine Skizze, weil es im Bildraum kein zweidimensionales Koordinatensystem mehr gibt. Mit einer Matrix wie

$$\begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}. \quad (24.2)$$

Projektion erhält man eine *Projektion*. Alle Punkte werden auf die horizontale Achse projiziert. Sie „verlieren“ quasi ihre zweite Koordinate. Anders ausgedrückt: Alle Punkte mit gleicher x -Koordinate werden auf denselben Punkt abgebildet.

Die gezeigten Abbildungstypen bilden einen „Baukasten“, aus dem man sich durch **Komposition** alle linearen Abbildungen zusammenbasteln kann. Das wird in den folgenden Aufgaben an diversen Beispielen durchexerziert.

Aufgabe 24.1. Geben Sie die Matrix für eine Skalierung an, die die gesamte Ebene um den Faktor 3.7 vergrößert.

Diagonalmatrix

Aufgabe 24.2. Eine Matrix, in der außerhalb der Hauptdiagonalen nur Nullen stehen, bezeichnet man als *Diagonalmatrix*. Offenbar entsprechen Diagonalmatrizen genau den allgemeinen Skalierungen, bei denen unabhängig voneinander in x - und y -Richtung skaliert wird. Überzeugen Sie sich durch Nachrechnen, dass das Produkt zweier Diagonalmatrizen eine Diagonalmatrix ist.

Aufgabe 24.3. Sei $\mathbf{A}$ eine $n \times n$ -Matrix und λ ein Skalar. Welche Matrix $\mathbf{M}$ sorgt dafür, dass $\mathbf{M} \cdot \mathbf{A} = \lambda \cdot \mathbf{A}$ gilt?

Aufgabe 24.4. Überlegen Sie sich, wie eine Skizze für die Scherung

$$\begin{pmatrix} 1 & 0 \\ -1.2 & 1 \end{pmatrix}$$

aussehen müsste.

Aufgabe 24.5. Es ist anschaulich naheliegend, dass die Komposition zweier Drehungen (um den Nullpunkt) eine Drehung ist, deren Drehwinkel sich durch die Addition der Drehwinkel der komponierten Abbildungen ergibt. Überzeugen Sie sich durch eine entsprechende Matrixmultiplikation und Anwendung der **Additionstheoreme**, dass das auch wirklich stimmt.

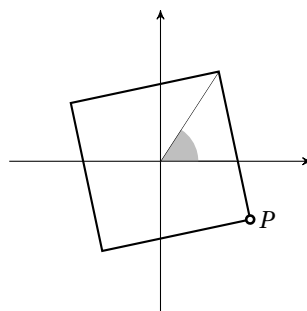
Aufgabe 24.6. Welche Bewegungen werden durch diese Matrizen beschrieben?

$$A = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix} \quad B = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$$

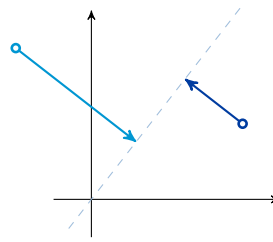
Aufgabe 24.7. Geben Sie die Drehmatrix für eine Drehung um 180° an und die Spiegelmatrix für eine Spiegelung an der x -Achse. Geben Sie diese Matrizen möglichst ohne Rechnung oder Verwendung trigonometrischer Funktionen an.

Aufgabe 24.8. Wenn man in der Ebene zweimal nacheinander um 180° dreht, dann erwartet man, dass das Ergebnis so ist, als hätte man gar keine Bewegung durchgeführt. Ebenso wird man erwarten können, dass eine zweimalige Spiegelung an derselben Achse (z. B. an der y -Achse) auch keinen Effekt hat. Überprüfen Sie das mit Matrizen. Stellen Sie also zunächst Matrizen für die Drehung bzw. die Spiegelung auf und berechnen Sie dann die entsprechenden Produkte.

Aufgabe 24.9. Das Quadrat in der folgenden Skizze hat den Mittelpunkt im Ursprung des Koordinatensystems und die Seitenlänge 2. Der grau hervorgehobene Winkel beträgt 57° . Wie berechnet man mit einer Drehmatrix die kartesischen Koordinaten des Eckpunktes P ?



Aufgabe 24.10. Geben Sie eine Matrix für die Abbildung an, die Punkte senkrecht auf die Gerade $y = 9x/7$ projiziert.



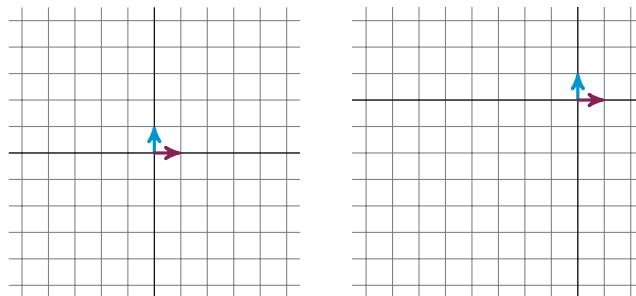
(Hinweis: Setzen Sie die Abbildung aus mehreren Teilen zusammen. Drehen Sie erst so, dass die Projektion sich einfacher als Projektion auf die horizontale Achse hinschreiben lässt, und drehen Sie anschließend wieder zurück. Setzen Sie wie in der Skizze testweise Punkte in Ihre Abbildung ein, um zu überprüfen, ob Ihr Ergebnis plausibel ist.)

Aufgabe 24.11. Geben Sie die Matrix für eine Spiegelung an der Geraden durch den Nullpunkt und den Punkt $(3, 1)$ an.

Aufgabe 24.12. Macht es einen Unterschied, ob man erst an der x -Achse und dann an der y -Achse spiegelt oder ob man es umgekehrt macht?

Translation

Eine wichtige Sorte von Funktionen fehlt in der obigen Liste. Es handelt sich um sogenannte *Translationen*, von denen als Beispiel $\mathbf{v} \mapsto \mathbf{v} + (3, 2)$ in der folgenden Skizze gezeigt wird. Man kann sofort sehen, dass es sich *nicht* um eine lineare Abbildung handeln kann, weil der Nullpunkt bewegt wird.



Außerdem geht durch Translationen der Zusammenhang zwischen Orts- und Richtungsvektoren verloren. Nennen wir die obige Beispielabbildung f , so gilt mit $\mathbf{p} = (2, 3)$, $\mathbf{q} = (3, 4)$ und $\mathbf{v} = (1, 1)$ zwar $\mathbf{q} = \mathbf{p} + \mathbf{v}$ und wir können $\mathbf{v}$ als Richtungsvektor auffassen, der den Punkt $\mathbf{p}$ nach $\mathbf{q}$ verschiebt, aber es gilt *nicht* $f(\mathbf{q}) = f(\mathbf{p}) + f(\mathbf{v})$ (rechnen Sie nach!), wie es bei einer linearen Abbildung der Fall wäre.

Dass man solche Abbildungen nicht durch eine Matrix darstellen kann, ist für die Anwendungen wohl das größte Problem, weil dadurch eine einheitliche Repräsentation aller relevanten Funktionen verhindert wird. Der nächste Abschnitt zeigt jedoch eine elegante Lösung.

25. Homogene Koordinaten

Das im letzten Abschnitt angesprochene Problem, dass Translationen keine linearen Abbildungen sind, lässt sich dadurch aus der Welt schaffen, dass man

gleichsam in eine höhere Dimension ausweicht. Formal arbeitet man mit sogenannten *homogenen Koordinaten*, die in der *projektiven Geometrie* verwendet werden.¹² Die für die Praxis wesentlichen Techniken kann man jedoch ohne die zugehörige Theorie recht leicht verstehen.

Dem Punkt $(a, b) \in \mathbb{R}^2$ ordnen wir seine **homogenen Koordinaten** $[a, b, 1] \in \mathbb{R}^3$ zu.¹³ Um den Unterschied hervorzuheben, verwenden wir in diesem Zusammenhang für Vektoren und Matrizen eckige Klammern. Gemeint sind jedoch dieselben Objekte wie vorher.

Definition

Man braucht nun lediglich zwei „Tricks“:

- Wird die lineare Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ durch die Matrix

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

dargestellt, dann hat die durch

$$\begin{bmatrix} a & b & 0 \\ c & d & 0 \\ 0 & 0 & 1 \end{bmatrix} \quad (25.1)$$

dargestellte lineare Abbildung $f^*: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ die folgende Eigenschaft: Bildet f den Vektor $\mathbf{v}$ auf den Vektor $\mathbf{w}$ ab, dann bildet f^* die homogenen Koordinaten von $\mathbf{v}$ auf die homogenen Koordinaten von $\mathbf{w}$ ab.

- Ist die Translation $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ durch $f(\mathbf{v}) = \mathbf{v} + (a, b)$ definiert, dann hat die durch

$$\begin{bmatrix} 1 & 0 & a \\ 0 & 1 & b \\ 0 & 0 & 1 \end{bmatrix} \quad (25.2)$$

dargestellte lineare Abbildung $f^*: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ die folgende Eigenschaft: Bildet f den Vektor $\mathbf{v}$ auf den Vektor $\mathbf{w}$ ab, dann bildet f^* die homogenen Koordinaten von $\mathbf{v}$ auf die homogenen Koordinaten von $\mathbf{w}$ ab.

Wie die beiden homogenen Matrizen entstehen, kann man sich leicht dadurch merken, dass die Einheitsmatrix $\mathbf{E}_3$ (graue Ziffern) teilweise überschrieben wird.

Aufgabe 25.1. Überzeugen Sie sich durch Nachrechnen, dass die beiden oben gezeigten „Tricks“ wirklich funktionieren.

Die geometrische Erklärung dafür, was hier passiert: Alle Punkte der Form $(x, y, 1)$ bilden im Raum $\mathbb{R}^3$ eine Ebene E , die parallel zur x - y -Ebene verläuft. Abbildungen

¹²Im Video *Wie muss man fragen, um "schöne" Antworten zu bekommen?* erfahren Sie mehr über projektive Geometrie.

¹³Das ist eine vereinfachte Definition, die für unsere Zwecke jedoch ausreicht.

vom Typ (25.1) oder (25.2) sind so gebaut, dass sie Punkte aus E auf Punkte aus E abbilden. (In der „dreidimensionalen Welt“ bilden sie Geraden durch den Nullpunkt auf Geraden durch den Nullpunkt ab. In der Ebene E „sieht“ man jeweils nur den Schnittpunkt dieser Geraden mit E .)

Matrizen vom Typ (25.1) übertragen im gewissen Sinne zweidimensionale lineare Abbildungen auf ihr dreidimensionales Pendant: Aus Drehungen um den Nullpunkt werden Drehungen um die z -Achse, aus Spiegelungen an Geraden werden Spiegelungen an Ebenen und so weiter. Matrizen vom Typ (25.2) beschreiben dreidimensionale Scherungen längs der x - y -Ebene. Innerhalb von E „erlebt“ man diese als Translationen.¹⁴

Wenn alle relevanten Abbildungen durch Matrizen repräsentiert werden können, dann können wir wie z. B. in Aufgabe 24.10 mehrere nacheinander ausgeführte Abbildungen zu einer Matrix zusammenfassen. Damit kann man Transformationen, die nicht linear sind, weil sie „am falschen Ort“ stattfinden, durch entsprechende Verschiebungen „linear machen“. Wir schauen uns das am Beispiel einer Drehung um $\pi/6$ um den Punkt $\mathbf{m} = (3, 2)$ an. Das ist keine lineare Abbildung, aber man kann sie in die folgenden drei Teile zerlegen:

- (i) Translation um $-\mathbf{m}$, so dass $\mathbf{m}$ auf den Ursprung verschoben wird
- (ii) die eigentliche Drehung, nun aber mit $\mathbf{0}$ als Mittelpunkt
- (iii) Translation um $\mathbf{m}$, um den ersten Schritt wieder rückgängig zu machen

Das Ergebnis sieht so aus:

$$\begin{aligned} & \begin{bmatrix} 1 & 0 & 3 \\ 0 & 1 & 2 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} \cos(\pi/6) & -\sin(\pi/6) & 0 \\ \sin(\pi/6) & \cos(\pi/6) & 0 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} 1 & 0 & -3 \\ 0 & 1 & -2 \\ 0 & 0 & 1 \end{bmatrix} \\ &= \begin{bmatrix} \sqrt{3}/2 & -1/2 & 4 - 3\sqrt{3}/2 \\ 1/2 & \sqrt{3}/2 & 1/2 - \sqrt{3} \\ 0 & 0 & 1 \end{bmatrix} \approx \begin{bmatrix} 0.865 & -0.500 & 1.402 \\ 0.500 & 0.865 & -1.232 \\ 0.000 & 0.000 & 1.000 \end{bmatrix} \end{aligned} \quad (25.3)$$

Aufgabe 25.2. Fällt Ihnen an (25.3) ein gewisses Einsparungspotential auf, wenn Sie das durch die Brille der Informatik betrachten?

Aufgabe 25.3. Sei f die Spiegelung an der y -Achse, g die durch $\mathbf{v} \mapsto \mathbf{v} + (3, 2)$ definierte Translation und h die Drehung um den Nullpunkt um 90° . (Damit sind natürlich drei Abbildungen von $\mathbb{R}^2$ nach $\mathbb{R}^2$ gemeint.) Geben Sie die homogene 3×3 -Matrix an, die die Abbildung $f \circ g \circ h$ repräsentiert.

Aufgabe 25.4. Stellen Sie die Spiegelung an der Geraden $g = \{(2, y) : y \in \mathbb{R}\}$ als Komposition von Translationen und linearen Abbildungen dar und generieren Sie daraus eine 3×3 -Matrix, die diese Spiegelung in homogenen Koordinaten realisiert.

¹⁴Ein eindimensionales Wesen, dass auf der Geraden $(0, 1) + \mathbb{R}(1, 0)$ lebt, würde die Scherung (24.1) ebenfalls als Translation wahrnehmen.

Mit homogenen Koordinaten kann man außer Translationen noch weitere Abbildungen von $\mathbb{R}^2$ nach $\mathbb{R}^2$, die nicht linear sind, durch 3×3 -Matrizen repräsentieren, z. B. die für perspektivische Darstellungen wichtige [Zentralprojektion](#). Und die hier vorgestellten Ideen funktionieren auch dann noch, wenn man im Raum $\mathbb{R}^3$ statt $\mathbb{R}^2$ arbeitet: Man verwendet dann 4×4 -Matrizen. In den Kapiteln 32 und 33 von [Konkrete Mathematik \(nicht nur\) für Informatiker](#) werden homogene Koordinaten sowie die Anwendung linearer Abbildungen für 3D-Darstellungen in der Computergrafik ausführlicher behandelt.

26. Inverse Matrizen

Ist $f: \mathbb{R}^m \rightarrow \mathbb{R}^n$ eine lineare Abbildung, die durch eine $n \times m$ -Matrix $\mathbf{A}$ beschrieben wird, so kann man die Frage, ob es einen Punkt $\mathbf{x} \in \mathbb{R}^m$ gibt, der von f auf den Punkt $\mathbf{b} \in \mathbb{R}^n$ abgebildet wird (für den also $f(\mathbf{x}) = \mathbf{b}$ gilt), als [lineares Gleichungssystem](#) $\mathbf{A} \cdot \mathbf{x} = \mathbf{b}$ auffassen. (Siehe Aufgabe 22.3.)

Im Spezialfall $m = n = 2$ stellt sich die Situation geometrisch wie folgt dar: Für Drehungen, Spiegelungen oder Scherungen ist es anschaulich evident, dass solche Abbildungen jeden Punkt der Ebene „treffen“. Das bedeutet, dass die zugehörigen linearen Gleichungssysteme unabhängig von der rechten Seite immer lösbar sind. Betrachtet man hingegen eine Projektion der Ebene auf eine Gerade, so ist klar, dass jeder Punkt, der nicht auf dieser Geraden liegt, nicht Bildpunkt der Abbildung ist. Daher kann das zugehörige lineare Gleichungssystem nicht lösbar sein, wenn die rechte Seite ein solcher Punkt ist.

Zurück zur allgemeinen Frage: Wenn obiges f injektiv wäre, wenn die [Umkehrfunktion](#) f^{-1} also existieren würde, dann könnten wir die in diesem Fall eindeutig bestimmte Lösung $\mathbf{x}$ von $\mathbf{A} \cdot \mathbf{x} = \mathbf{b}$ sofort als $\mathbf{x} = f^{-1}(\mathbf{b})$ berechnen. Wann gibt es also so eine Umkehrfunktion? Und wie sieht sie aus?

Diese Frage lässt sich wieder geometrisch beantworten. Wenn f eine Gerade auf einen Punkt abbildet, dann ist f sicherlich *nicht* injektiv, also gibt es keine Umkehrfunktion. Wenn f jedoch alle Geraden auf Geraden abbildet und die anderen Bedingungen von Lemma 23.1 erfüllt, dann muss es eine Umkehrfunktion geben. Außerdem kann man sich leicht überlegen, dass diese Umkehrfunktion natürlich auch wieder Geraden auf Geraden abbilden muss und dass sie ebenfalls die Streckenverhältnisse und die Parallelität bewahren muss. (Das ist kein mathematischer Beweis, aber anschaulich sollte klar sein, dass es so funktionieren muss.)

Wenn es eine Umkehrfunktion gibt, dann ist sie also nicht irgendeine Funktion, sondern auch eine lineare Abbildung und muss sich daher auch durch eine Matrix darstellen lassen. Bevor wir uns aber nun Gedanken darüber machen, wie man diese Matrix berechnet, müssen wir noch kurz überlegen, für welche Typen von Matrizen so eine Berechnung überhaupt einen Sinn ergibt.

Haben wir es mit einer Matrix zu tun, die mehr Spalten als Zeilen hat, so wird ein lineares Gleichungssystem zu dieser Koeffizientenmatrix, wenn es überhaupt lösbar ist, mehr als eine Lösung haben. (Stellen Sie sich dazu die Zeilenstufenform

einer solchen Matrix vor.) Das bedeutet aber, dass die zugehörige lineare Abbildung mit Sicherheit *nicht* injektiv ist. Hat die Koeffizientenmatrix umgekehrt mehr Zeilen als Spalten, so hat sie nach Umwandlung in Zeilenstufenform Nullzeilen. Das impliziert jedoch, dass das System nicht für *alle* rechten Seiten lösbar sein kann. Die Umkehrfunktion kann somit keine lineare Abbildung sein, die auf der gesamten Zielmenge $\mathbb{R}^n$ definiert ist.

Die Konsequenz dieser Überlegungen ist, dass wir uns auf *quadratische* Matrizen beschränken können. Sei also $\mathbf{A} \in \mathbb{R}^{n \times n}$ und f die zugehörige lineare Abbildung:

$$f: \begin{cases} \mathbb{R}^n \rightarrow \mathbb{R}^n \\ \mathbf{x} \mapsto \mathbf{A} \cdot \mathbf{x} \end{cases}$$

Wir suchen die Matrix $\mathbf{B}$, die die Umkehrabbildung f^{-1} darstellt. Nach Satz 16.3 gilt $f(f^{-1}(\mathbf{x})) = \mathbf{x}$ und $f^{-1}(f(\mathbf{x})) = \mathbf{x}$ für alle $\mathbf{x} \in \mathbb{R}^n$. Außerdem gilt für solche $\mathbf{x}$ immer $\mathbf{E}_n \mathbf{x} = \mathbf{x}$. Daher ergibt die folgende Definition Sinn:

Definition

Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$. Gibt es eine Matrix $\mathbf{B} \in \mathbb{R}^{n \times n}$ mit

$$\mathbf{A} \cdot \mathbf{B} = \mathbf{E}_n = \mathbf{B} \cdot \mathbf{A},$$

dann nennen wir sie die **inverse Matrix**¹⁵ zu $\mathbf{A}$ und schreiben sie als $\mathbf{A}^{-1}$. Matrizen, zu denen es inverse Matrizen gibt, nennt man **invertierbar** oder auch **regulär**. Eine Matrix, die nicht regulär ist, wird **singulär** genannt.

Aufgabe 26.1. Geben Sie eine 2×2 -Matrix an, die nicht regulär ist.

Haben wir es mit einer invertierbaren Matrix $\mathbf{A}$ zu tun, so ist jedes lineare Gleichungssystem, dessen Koeffizientenmatrix $\mathbf{A}$ ist, eindeutig lösbar. Daher muss sich bei der Anwendung des **Gauß-Jordan-Verfahrens** auf $\mathbf{A}$ die Einheitsmatrix ergeben. Am Ende von Abschnitt 22 haben wir gelernt, dass sich elementare Zeilenumformungen als Multiplikationen mit bestimmten Matrizen darstellen lassen. Der Prozess, der aus $\mathbf{A}$ die Einheitsmatrix macht, sieht also folgendermaßen aus:

$$\mathbf{D}_k \cdots \mathbf{D}_2 \cdot \mathbf{D}_1 \cdot \mathbf{A} = \mathbf{E}_n \quad (26.1)$$

Das ist so gemeint, dass wir k elementare Zeilenumformungen durchgeführt haben. Der ersten Umformung entspricht die Matrix $\mathbf{D}_1$, der zweiten die Matrix $\mathbf{D}_2$, und so weiter. Da die Inverse $\mathbf{A}^{-1}$ existiert, können wir beide Seiten von (26.1) mit dieser Matrix multiplizieren,¹⁶ ohne die Gültigkeit der Gleichung zu verletzen:

$$\mathbf{D}_k \cdots \mathbf{D}_2 \cdot \mathbf{D}_1 \cdot \mathbf{A} \cdot \mathbf{A}^{-1} = \mathbf{E}_n \cdot \mathbf{A}^{-1}$$

¹⁵Der Begriff ist sicher sinnvoll, da es sich ja um das **inverse Element** zu $\mathbf{A}$ bzgl. der Matrizenmultiplikation handelt.

¹⁶Da die Matrizenmultiplikation nicht kommutativ ist, müssen wir entweder auf beiden Seiten von links mit derselben Matrix multiplizieren oder wie hier auf beiden Seiten von rechts.

Da einerseits $\mathbf{A} \cdot \mathbf{A}^{-1} = \mathbf{E}_n$ gilt und andererseits die Einheitsmatrix das neutrale Element der Matrizenmultiplikation ist, vereinfacht sich das:

$$\mathbf{D}_k \cdots \mathbf{D}_2 \cdot \mathbf{D}_1 \cdot \mathbf{E}_n = \mathbf{A}^{-1}$$

Das bedeutet, dass wir nur dieselben Zeilenumformungen in derselben Reihenfolge auf $\mathbf{E}_n$ anwenden müssen, um $\mathbf{A}^{-1}$ zu erhalten! Wir erhalten so ein Verfahren zum Invertieren von Matrizen:

- Man schreibe die zu invertierende Matrix $\mathbf{A}$ und die Einheitsmatrix vom selben Typ nebeneinander.
- Wie im Gauß-Jordan-Verfahren wird aus $\mathbf{A}$ nun durch elementare Zeilenumformungen die Einheitsmatrix gemacht.
- Jeder Schritt, der an $\mathbf{A}$ durchgeführt wird, wird parallel auch an der Einheitsmatrix vorgenommen.
- Am Ende ist aus der Einheitsmatrix $\mathbf{A}^{-1}$ geworden.
- Kommt man zwischendurch an einen Punkt, an dem es nicht weitergeht, so ist $\mathbf{A}$ nicht regulär und man kann aufhören.¹⁷ Dieses Verfahren ist also gleichzeitig ein Test, um festzustellen, ob eine Matrix invertierbar ist.

Hier ein Beispiel. Die zu invertierende Matrix steht links oben, daneben $\mathbf{E}_3$. Die Inverse steht rechts unten. Die Pivotelemente sind rot eingefärbt.

$$\begin{pmatrix} 1 & 0 & 2 \\ 2 & -1 & 3 \\ 4 & 1 & 8 \end{pmatrix} \quad \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

$$\begin{pmatrix} 1 & 0 & 2 \\ 0 & -1 & -1 \\ 0 & 1 & 0 \end{pmatrix} \quad \begin{pmatrix} 1 & 0 & 0 \\ -2 & 1 & 0 \\ -4 & 0 & 1 \end{pmatrix}$$

$$\begin{pmatrix} 1 & 0 & 2 \\ 0 & -1 & -1 \\ 0 & 0 & -1 \end{pmatrix} \quad \begin{pmatrix} 1 & 0 & 0 \\ -2 & 1 & 0 \\ -6 & 1 & 1 \end{pmatrix}$$

$$\begin{pmatrix} 1 & 0 & 2 \\ 0 & -1 & -1 \\ 0 & 0 & 1 \end{pmatrix} \quad \begin{pmatrix} 1 & 0 & 0 \\ -2 & 1 & 0 \\ 6 & -1 & -1 \end{pmatrix}$$

$$\begin{pmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad \begin{pmatrix} -11 & 2 & 2 \\ 4 & 0 & -1 \\ 6 & -1 & -1 \end{pmatrix}$$

$$\begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad \begin{pmatrix} -11 & 2 & 2 \\ -4 & 0 & 1 \\ 6 & -1 & -1 \end{pmatrix}$$

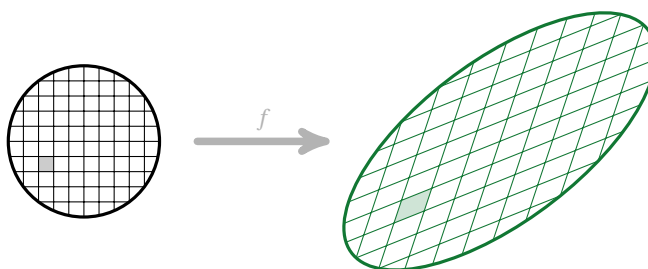
¹⁷Weiß man also noch nicht, ob $\mathbf{A}$ regulär ist, so ist es effizienter, die Umformungen an der Einheitsmatrix erst ganz zum Schluss durchzuführen, da man sonst evtl. unnötige Arbeitsschritte ausführt.

Aufgabe 26.2. Verifizieren Sie, dass die beiden Matrizen wirklich invers zueinander sind. (Wenn man die Matrix links oben mit der rechts unten multipliziert – egal in welcher Reihenfolge –, muss die Einheitsmatrix herauskommen.)

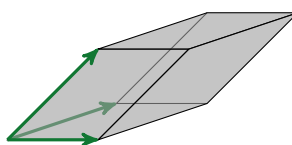
Aufgabe 26.3. Üben Sie das Invertieren von Matrizen. (Es ist offensichtlich einfach, sich selbst Aufgaben zu stellen, weil man das Ergebnis wie in Aufgabe 26.2 sofort überprüfen kann.)

27. Determinanten

In Abschnitt 24 haben wir gesehen, dass eine lineare Abbildung f von $\mathbb{R}^2$ nach $\mathbb{R}^2$ aus dem Netz identischer Quadrate, das die Ebene überdeckt, ein Netz identischer Parallelogramme macht. Die Flächeninhalte der Quadrate werden also durch f alle um denselben Faktor vergrößert oder verkleinert. Und offensichtlich gilt das nicht nur für die speziellen Gitterquadrate, sondern für *alle* Quadrate. Es gilt sogar für *alle Flächen*, weil wir uns die wie in der folgenden Skizze aus sehr kleinen Quadraten zusammengesetzt vorstellen können.¹⁸



Zur linearen Abbildung f gehört also offenbar ein fester „Flächenfaktor“, also eine Zahl. Und dieselbe Überlegung kann man auch im dreidimensionalen Raum anstellen. Dort werden durch lineare Abbildungen aus **Würfeln** sogenannte *Spat* bzw. *Parallelepipede* gemacht:



In diesem Fall kann man sich mit analogen Überlegungen klarmachen, dass zu jeder linearen Abbildung von $\mathbb{R}^3$ nach $\mathbb{R}^3$ eine feste Zahl gehört, die angibt, wie durch diese Abbildung das Volumen beeinflusst wird.

Und diese Zahl kann man in beiden Fällen¹⁹ anhand der Matrix $\mathbf{A}$, die f beschreibt berechnen. Man nennt sie die *Determinante* der Matrix $\mathbf{A}$ und schreibt dafür $\det(\mathbf{A})$.

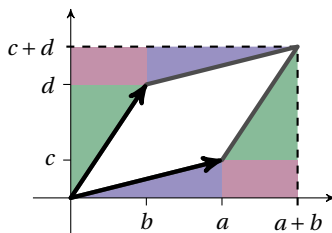
¹⁸Das ist natürlich nur ein anschauliches Argument. Man kann das jedoch mithilfe der **Integralrechnung** präzisieren.

¹⁹Das gilt sogar für höhere Dimensionen, die man sich bildlich nicht mehr vorstellen kann. Man spricht dann vom *verallgemeinerten Volumen*. In diesem Sinne ist der Flächeninhalt das zweidimensionale verallgemeinerte Volumen.

Wir werden im Folgenden entwickeln, wie man diese Determinante berechnen kann und dass sie noch weitere Informationen bereitstellt. Warum das im Detail möglich ist, werden wir nicht immer ausführlich begründen können, weil das zu weit führen würde. Teilweise kann man sich jedoch zumindest für den zweidimensionalen Fall mit Skizzen von der Korrektheit der Behauptungen überzeugen.

Zunächst die zweidimensionale Situation. Von dem folgenden weißen Parallelogramm sagt man, dass es von den durch Pfeile ange deuteten Vektoren (a, c) und (b, d) *aufgespannt* wird.

aufspannen



Aufgabe 27.1. Berechnen Sie die Fläche des Parallelogramms, indem Sie von der Fläche großen gestrichelten Rechtecks, die Flächen der farbig unterlegten kleinen Rechtecke bzw. rechtwinkligen Dreiecke abziehen.

Aufgabe 27.1 zeigt, dass man die Determinante einer 2×2 -Matrix nach dem Schema

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

berechnen kann. Wir werden wie allgemein üblich für Determinanten in Zukunft auch die Schreibweise

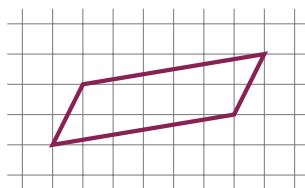
$$\begin{vmatrix} a & b \\ c & d \end{vmatrix} = ad - bc \quad (27.1)$$

mit vertikalen Strichen verwenden.

Aufgabe 27.2. Berechnen Sie diese beiden Determinanten:

$$\begin{vmatrix} 5 & 2 \\ -8 & -3 \end{vmatrix} \quad \begin{vmatrix} 42 & 1/3 \\ 6 & 0 \end{vmatrix}$$

Aufgabe 27.3. Welche Fläche hat das folgende Parallelogramm? (Dabei sollen benachbarte parallele Gitterlinien den Abstand 1 haben.)



Aufgabe 27.4. Geben Sie die Fläche des Parallelogramms mit den Eckpunkten $(2, 1)$, $(7, 4)$ und $(3, 5)$ an. (Nein, es wurde kein Eckpunkt vergessen. Drei Ecken determinieren ein Dreieck, dessen Fläche die Hälfte der Parallelogrammfläche ausmacht. Damit gibt es für die Antwort nur eine Möglichkeit.)

Aufgabe 27.5. Sei $\mathbf{b}$ der Vektor $(3, 2)$. Geben Sie einen zur horizontalen Achse parallelen Vektor $\mathbf{a}$ an, für den das von $\mathbf{a}$ und $\mathbf{b}$ aufgespannte Parallelogramm die Fläche 5 hat.

Aufgabe 27.6. Bei einer Matrix wie

$$\begin{vmatrix} 3 & 6 \\ 1 & 2 \end{vmatrix},$$

bei der eine Spalte das skalare Vielfache einer anderen ist, ergibt sich als Determinante null. (Rechnen Sie nach!) Was ist der geometrische Grund dafür?

Aufgabe 27.7. Welche Determinante ergibt sich, wenn man die Matrix in (27.1) transponiert?

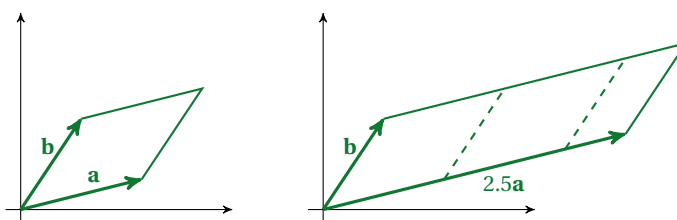
Aufgabe 27.8. Sie sollten aus der Schule noch wissen, dass man die Fläche eines Parallelogramms mittels „Grundseite mal Höhe“ berechnen kann. Machen Sie sich anhand von Skizzen klar, warum das funktioniert, indem Sie von Parallelogrammen geeignete Dreiecke „abschneiden“.

In Aufgabe 27.2 ist Ihnen vielleicht schon aufgefallen, dass sich beim Berechnen der Determinante negative Zahlen ergeben können. Wie passt das zu der Aussage, dass es um Flächeninhalte geht? Die Antwort ist, dass die Determinante die *vorzeichenbehaftete* Fläche angibt. Ihr **Absolutbetrag** gibt die Fläche an, ihr Vorzeichen die Orientierung. Die Matrizen

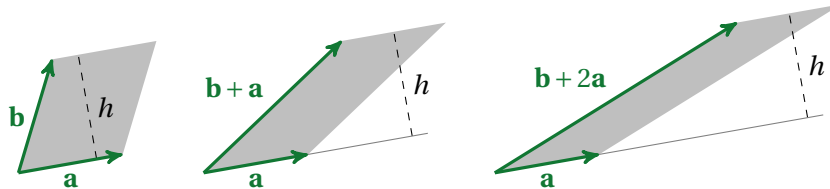
$$\begin{pmatrix} a & b \\ c & d \end{pmatrix} \quad \text{und} \quad \begin{pmatrix} b & a \\ d & c \end{pmatrix},$$

bei denen die Spalten vertauscht wurden, führen zu zwei Parallelogrammen, die von denselben beiden Vektoren (a, c) und (b, d) aufgespannt werden, aber die Drehrichtung vom ersten auf den zweiten Vektor ist unterschiedlich. Im ersten Fall ist die Determinante $ad - bc$, im zweiten $bc - ad = -(ad - bc)$.

Bei einer weiteren Änderung, die man an der Matrix vornehmen kann, sollte anschaulich sofort klar sein, wie sie sich auf die Determinante auswirkt: Multipliziert man einen der Vektoren mit einem Skalar λ , so wird die Fläche offenbar um den Faktor λ vergrößert bzw. verkleinert.



Ein letzter Baustein fehlt noch. Das von den Vektoren $\mathbf{a}$ und $\mathbf{b}$ aufgespannte Parallelogramm hat dieselbe Fläche wie das, das von den Vektoren $\mathbf{a}$ und $\mathbf{b} + \lambda \mathbf{a}$ aufgespannt wird, wobei λ irgendein Skalar sein darf. Das folgt aus der bekannten Flächenformel, um die es in Aufgabe 27.8 ging, weil sich die senkrecht auf $\mathbf{a}$ stehende Höhe nicht ändert. Die folgende Skizze zeigt dazu zwei Beispiele für $\lambda = 1$ und $\lambda = 2$.

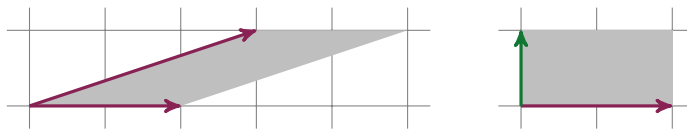


Schließlich gibt es noch Matrizen, deren Determinanten sich besonders einfach ausrechnen lassen, z. B. diese:

$$\mathbf{A} = \begin{pmatrix} a & b \\ 0 & d \end{pmatrix}$$

Verwendet man (27.1), so sieht man, dass der Wert $\det(\mathbf{A})$ sich einfach als das Produkt ad der Einträge auf der Hauptdiagonalen ergibt, was offenbar an der Null liegt.

Auch geometrisch kann man sich das veranschaulichen.



Da der erste Vektor parallel zur horizontalen Achse verläuft, ist seine Länge seine erste Komponente. Zur Berechnung der Fläche brauchen wir dann lediglich noch die Höhe auf diesem Vektor. Das ist aber einfach die zweite Komponente des zweiten Vektors. Die Fläche würde sich nicht ändern, wenn die andere Komponente des zweiten Vektors andere Werte hätten.

Wir fassen nun die bisherigen Überlegungen zusammen. Sie gelten alle auch für höhere Dimensionen und werden entsprechend formuliert. Die Begründungen dafür sind allerdings teilweise komplizierter.

Jeder quadratischen Matrix $\mathbf{A}$ lässt sich ein eindeutig bestimmter Skalar $\det(\mathbf{A})$ zuordnen, den man die *Determinante* der Matrix nennt. Die Determinante hat die folgenden Eigenschaften.

- Determinanten von Matrizen in oberer Dreiecksform²⁰ ergeben sich als Produkt der Einträge in der Hauptdiagonalen.
- Vertauscht man zwei Spalten einer Matrix, so ändert sich dadurch das Vorzeichen der Determinante.

Satz 27.1

- Multipliziert man eine Spalte mit einem Faktor, dann ist auch die sich so ergebende Determinante um diesen Faktor größer bzw. kleiner als die ursprüngliche.
- Addiert man zu einer Spalte das Vielfache einer anderen hinzu, so ändert sich dadurch die Determinante nicht.
- Die Determinante ändert sich durch Transponieren nicht, d. h., es gilt $\det(\mathbf{A}) = \det(\mathbf{A}^T)$. Daher gelten die letzten drei Aussagen auch für Zeilen statt für Spalten.

Leider gibt es für größere Matrizen keine so einfachen Formeln wie (27.1). Durch das Hinzufügen weiterer Zeilen bzw. Spalten steigt der Rechenaufwand erheblich an. Die einfachste Möglichkeit, Determinanten für beliebige Matrizen zu berechnen, ist das Gaußverfahren, das wir in Abschnitt 22 kennengelernt haben. Damit kann man eine Matrix auf obere Dreiecksform bringen, wodurch ihre Determinante nach Satz 27.1 leicht zu berechnen ist. Und dieser Satz sagt auch, welche Auswirkungen die im Laufe des Verfahrens angewendeten elementaren Zeilenumformungen auf das Ergebnis haben.

Man kann zum Berechnen der Determinante einer beliebigen Matrix daher so vorgehen: Man bringt die Matrix durch elementare Zeilenumformungen auf obere Dreiecksform und führt dabei Buch über die Operationen, die den Wert der Determinante ändern. Die Determinante der Dreiecksmatrix lässt sich sofort berechnen. Aus diesem Wert erhält man die gesuchte Determinante, indem man die notierten Änderungen wieder rückgängig macht.

Das klingt vielleicht komplizierter, als es ist. Wir schauen uns das am Beispiel der elementaren Zeilenumformungen von Aufgabe 22.6 an:

$$\begin{aligned}
 \mathbf{A} &= \begin{pmatrix} 6 & -1 & -1 \\ 1 & 1 & 10 \\ 2 & -1 & 1 \end{pmatrix} \\
 &\begin{pmatrix} \textcircled{1} & 1 & 10 \\ 6 & -1 & -1 \\ 2 & -1 & 1 \end{pmatrix} \quad \boxed{\cdot(-1)} \\
 &\begin{pmatrix} \textcircled{1} & 1 & 10 \\ 0 & -7 & -61 \\ 2 & -1 & 1 \end{pmatrix} \\
 &\begin{pmatrix} 1 & 1 & 10 \\ 0 & -7 & -61 \\ 0 & -3 & -19 \end{pmatrix} \\
 &\begin{pmatrix} 1 & 1 & 10 \\ 0 & \textcircled{-7} & -61 \\ 0 & -21 & -133 \end{pmatrix} \quad \boxed{\cdot 7}
 \end{aligned}$$

²⁰Siehe Aufgabe 22.5.

$$\mathbf{A}' = \begin{pmatrix} 1 & 1 & 10 \\ 0 & -7 & -61 \\ 0 & 0 & 50 \end{pmatrix}$$

Am Ende haben wir eine Matrix, deren Determinante sich ganz einfach berechnen lässt: $\det(\mathbf{A}') = 1 \cdot (-7) \cdot 50 = -350$. Allerdings war unsere erste Zeilenumformung die Vertauschung zweier Zeilen. Das entspricht einer Multiplikation mit -1 , was wir am Rande vermerkt haben. Im vorletzten Schritt haben wir eine Zeile mit 7 multipliziert und dies ebenfalls am Rande notiert. Insgesamt haben wir also mit $(-1) \cdot 7 = -7$ multipliziert und müssen das am Ende wieder korrigieren. Daher ist $\det(\mathbf{A}) = 1/(-7) \cdot \det(\mathbf{A}') = 50$.

Aufgabe 27.9. Berechnen Sie die Determinanten der beiden folgenden Matrizen:

$$\mathbf{A} = \begin{pmatrix} 1 & 1 & 2 & -4 \\ 2 & 1 & 5 & 0 \\ 0 & 2 & -1 & -1 \\ 7 & 0 & 2 & 3 \end{pmatrix} \quad \mathbf{B} = \begin{pmatrix} -1 & 3 & 4 & 1 \\ 3 & 1 & 0 & 5 \\ -3 & -2 & 0 & 3 \\ -2 & 1 & 1 & 0 \end{pmatrix}$$

Für 3×3 -Matrizen gibt es eine Formel, die manchmal etwas einfacher anzuwenden ist als das Gaußverfahren. Die Determinante der Matrix

$$\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix}$$

ist

$$\begin{aligned} & a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} \\ & - a_{13}a_{22}a_{31} - a_{12}a_{21}a_{33} - a_{11}a_{23}a_{32} \end{aligned} \quad (27.2)$$

Für (27.2) gibt es eine grafische Eselsbrücke:

Man wiederholt neben der Matrix im Geiste ihre ersten beiden Spalten. Das wird auch *Regel von Sarrus* genannt.

Regel von Sarrus

Aufgabe 27.10. Berechnen Sie die folgenden Determinanten:

$$\begin{vmatrix} 1 & 3 & 5 \\ 8 & 7 & 0 \\ 2 & 1 & -4 \end{vmatrix} \quad \begin{vmatrix} 1 & 3 & 5 \\ 1 & -2 & -9 \\ 2 & 1 & -4 \end{vmatrix}$$

Aufgabe 27.11. Berechnen Sie die folgende Determinante im Restklassenkörper $\mathbb{Z}_{11}$.

$$\begin{vmatrix} 4 & 5 & 8 \\ 1 & 0 & 8 \\ 3 & 0 & 4 \end{vmatrix}$$

Aufgabe 27.12. Unten sind vier Matrizen aufgeführt, die wir so oder so ähnlich im Abschnitt 24 kennengelernt haben. Geben Sie jeweils an, welche geometrische Transformation beschrieben wird. Berechnen Sie dann die Determinanten und interpretieren Sie die Werte geometrisch.

$$\mathbf{A} = \begin{pmatrix} \cos(\pi/5) & -\sin(\pi/5) \\ \sin(\pi/5) & \cos(\pi/5) \end{pmatrix} \quad \mathbf{B} = \begin{pmatrix} \cos(\pi/3) & \sin(\pi/3) \\ \sin(\pi/3) & -\cos(\pi/3) \end{pmatrix}$$

$$\mathbf{C} = \begin{pmatrix} 1 & 2.5 \\ 0 & 1 \end{pmatrix} \quad \mathbf{D} = \begin{pmatrix} 3 & 0 \\ 0 & 1 \end{pmatrix}$$

Aufgabe 27.13. Welche Determinante hat eine Einheitsmatrix?

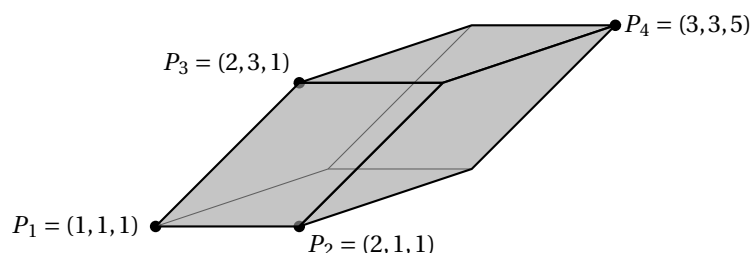
Aufgabe 27.14. Geben Sie eine 2×2 -Matrix $\mathbf{A}$ an, für die $\det(-\mathbf{A}) \neq -\det(\mathbf{A})$ gilt.

Aufgabe 27.15. Sei $\mathbf{A}$ eine $n \times n$ -Matrix. Wie kann man den Wert $\det(\lambda \cdot \mathbf{A})$ in Abhängigkeit von $\det(\mathbf{A})$ ausdrücken? Entspricht das Ihrer geometrischen Intuition?

Aufgabe 27.16. Geben Sie die Determinanten von $-\mathbf{E}_{21}$ und von $\mathbf{E}_{10} + \mathbf{E}_{10}$ an.

Aufgabe 27.17. Die durch $(v_1, v_2) \mapsto (v_1 + 7v_2, v_1 + 3v_2)$ definierte lineare Abbildung bildet den Einheitskreis auf eine Ellipse ab. Welche Fläche hat diese Ellipse?

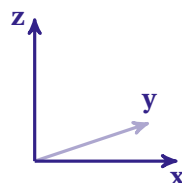
Aufgabe 27.18. Berechnen Sie das Volumen des folgenden Spats:



(Der Spat wurde *nicht* maßstabsgetreu gezeichnet. Relevant sind nur die Koordinaten.)

Drei-Finger-Regel

Wir hatten weiter oben bereits über die geometrische Bedeutung des Vorzeichens bei 2×2 -Matrizen gesprochen. Wie ist das im dreidimensionalen Raum? Dort ergibt sich das Vorzeichen analog zur aus der Schule bekannten *Drei-Finger-Regel*: Sind die drei Vektoren, die einen Spat aufspannen, in der Reihenfolge angeordnet, die den Richtungen von Daumen, Zeige- und Mittelfinger der rechten Hand entspricht, so ist die Determinante positiv, anderenfalls negativ. Wenn wir uns in der folgenden Skizze z. B. vorstellen, dass $\mathbf{x}$ und $\mathbf{z}$ in der Bildebene liegen und $\mathbf{y}$ von uns weg zeigt, so entspricht der Reihenfolge $\mathbf{x}, \mathbf{y}, \mathbf{z}$ eine positive Determinante, $\mathbf{x}, \mathbf{z}, \mathbf{y}$ aber z. B. eine negative.



Aufgabe 27.19. Machen Sie sich – vielleicht durch ein paar Skizzen oder mit drei Strohhalmen – klar, dass es tatsächlich nur zwei Möglichkeiten für den Drehsinn gibt, obwohl man drei Vektoren auf sechs verschiedene Arten sortieren kann. Und dass das Vertauschen zweier Vektoren immer zu einem Wechsel des Drehsinns führt.

Wenn man Determinanten als Flächen- bzw. Volumenfaktoren von linearen Abbildungen betrachtet und sich **daran erinnert**, dass die Komposition solcher Funktionen der Multiplikation der Matrizen entspricht, ist die folgende Aussage sofort klar:

Sind **A** und **B** zwei quadratische Matrizen desselben Typs, so gilt für ihr Produkt $\det(\mathbf{A} \cdot \mathbf{B}) = \det(\mathbf{A}) \cdot \det(\mathbf{B})$.

Lemma 27.2

Aufgabe 27.20. Gegeben sei die folgende Matrix **A**. Berechnen Sie die Determinante von $\mathbf{A}^5 = \mathbf{A} \cdot \mathbf{A} \cdot \mathbf{A} \cdot \mathbf{A} \cdot \mathbf{A}$.

$$\mathbf{A} = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & -6 \\ 2 & 3 & 1 \end{pmatrix}$$

Aufgabe 27.21. Wenn die Matrix **A** **regulär** ist, wie kann man dann die Determinante von $\mathbf{A}^{-1}$ ganz einfach ausrechnen, wenn man die von **A** kennt? (Hinweis: Was ist $\mathbf{A} \cdot \mathbf{A}^{-1}$? Wenden Sie Lemma 27.2 an.)

Aufgabe 27.22. Gegeben sei die folgende Matrix:

$$\mathbf{A} = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 6 & 5 \\ 3 & 2 & 1 \end{pmatrix}$$

Geben Sie die Determinante von $\mathbf{A}^{-1}$ an.

Aufgabe 27.23. Die folgende Aussage gilt im Allgemeinen übrigens *nicht*:

$$\det(\mathbf{A} + \mathbf{B}) = \det(\mathbf{A}) + \det(\mathbf{B})$$

Geben Sie ein konkretes Gegenbeispiel an. (Hinweis: Man kommt mit sehr einfachen 2×2 -Matrizen aus. Probieren Sie einfach ein bisschen herum.)

Aufgabe 27.24. Geben Sie eine 2×2 -Matrix an, für die $\det(\mathbf{A} + \mathbf{A}) = \det(\mathbf{A}) + \det(\mathbf{A})$ gilt. Die Nullmatrix zählt nicht als Antwort.

Aufgabe 27.25. Welche beiden Werte für a sorgen dafür, dass die Matrix $\mathbf{A} \cdot \mathbf{A}$ die Determinante 9 hat?

$$\mathbf{A} = \begin{pmatrix} 1 & 1 & -1 \\ 3 & a & 0 \\ 1 & 1 & 2 \end{pmatrix}$$

Die folgenden Aussage stellt Zusammenhänge zwischen den Konzepten aus diesem Abschnitt und denen davor her. Einige dieser Zusammenhänge sollte klar sein, bei einigen anderen müssen wir uns ohne Beweis auf sie verlassen. Es handelt sich jedenfalls um einen der zentralen Sätze der linearen Algebra.

Satz 27.3

Ist $\mathbf{A}$ eine $n \times n$ -Matrix und f die durch $\mathbf{v} \mapsto \mathbf{A}\mathbf{v}$ definierte lineare Abbildung, so sind die folgenden Aussagen alle äquivalent:²¹

- (i) $\mathbf{A}$ ist regulär.
- (ii) $\mathbf{A}^T$ ist regulär.
- (iii) Die Determinante von $\mathbf{A}$ ist *nicht* null.
- (iv) f ist injektiv.
- (v) Der Wertebereich von f ist $\mathbb{R}^n$.
- (vi) f ist injektiv und der Wertebereich von f ist $\mathbb{R}^n$.
- (vii) Der Nullvektor ist die einzige Lösung des LGS $\mathbf{A}\mathbf{x} = \mathbf{0}$.
- (viii) Das LGS $\mathbf{A}\mathbf{x} = \mathbf{b}$ hat für jede rechte Seite $\mathbf{b}$ genau eine Lösung.

Aufgabe 27.26. Geben Sie eine 2×2 -Matrix an, die singulär ist und von deren vier Komponenten mindestens zwei nicht null sind. (Tipp: Nach Satz 27.3 lässt sich für so kleine Matrizen Singularität am einfachsten mithilfe der Determinante testen.)

Aufgabe 27.27. Geben Sie eine 2×2 -Matrix an, die singulär ist und deren vier Komponenten alle verschieden sind.

Aufgabe 27.28. Ist die Funktion $(x, y) \mapsto (3x - 2y, 2x + y)$ von $\mathbb{R}^2$ nach $\mathbb{R}^2$ injektiv?

Aufgabe 27.29. Ist die Funktion $(x, y) \mapsto (3x - 2y, 2x + 1)$ von $\mathbb{R}^2$ nach $\mathbb{R}^2$ injektiv? (Hinweis: Die Funktion ist keine lineare Abbildung, lässt sich aber als Komposition zweier Funktionen darstellen.)

Aufgabe 27.30. Welche der folgenden Aussagen sind für alle Matrizen $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times n}$ wahr? Geben Sie bei den Aussagen, die nicht wahr sind, jeweils ein konkretes Gegenbeispiel an.

- (i) $\det(\mathbf{A}^{-1}) = \det(\mathbf{A})$
- (ii) $\det(\mathbf{A} \cdot \mathbf{B}) = \det(\mathbf{A}) \cdot \det(\mathbf{B})$
- (iii) $\det(\mathbf{A}^T) = \det(\mathbf{A})$
- (iv) $\det(\mathbf{A} + \mathbf{B}) = \det(\mathbf{A}) + \det(\mathbf{B})$
- (v) $\det(42 \cdot \mathbf{A}) = 42 \cdot \det(\mathbf{A})$
- (vi) $\det(\mathbf{A} \cdot \mathbf{A}) = \det(\mathbf{A})^2$
- (vii) $\det(\mathbf{A}^{-1}) = \det(\mathbf{A}^T)$
- (viii) $\det(-\mathbf{A}) = -\det(\mathbf{A})$

²¹Sie wissen aus Kapitel I, was das bedeutet: Wenn eine der Aussagen gilt, dann gelten auch alle anderen. Gilt eine von ihnen nicht, so gilt auch keine der anderen.

Bemerkungen

Beachten Sie den Unterschied zwischen

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix} \quad \text{und} \quad \begin{vmatrix} a & b \\ c & d \end{vmatrix}.$$

Links steht eine Matrix, rechts deren Determinante, also eine Zahl!

Für die Determinante der Matrix **A** schreibt man jedoch $\det(\mathbf{A})$ und *nicht* $|\mathbf{A}|$.

Ein häufiger Anfängerfehler ist, die Regel von Sarrus auf Matrizen anderen Typs zu übertragen. Diese Regel gilt *ausschließlich* für 3×3 -Matrizen!

★ Determinanten im Computer

Mit der [bereits definierten](#) Funktion `rowEchelon` kann man die Determinante einer Matrix ganz einfach berechnen:

```
def determinant (M):
    rowEchelon(M)
    det = 1
    for i in range(len(M)):
        det *= M[i][i]
    return det
```

Vergessen Sie aber nicht die Anmerkungen zu `rowEchelon` [im entsprechenden Abschnitt!](#)

28. Skalar- und Vektorprodukt

Bisher konnten wir Vektoren nur addieren und skalieren. In diesem Kapitel wird es nun um zwei ganz unterschiedliche Arten gehen, Vektoren zu *multiplizieren*. Obwohl man insbesondere die erste Art der Multiplikation ganz simpel algebraisch über die Komponenten definieren kann, wird sich herausstellen, dass sie eine koordinatenunabhängige geometrische Bedeutung hat und uns eine neue Sichtweise auf Längen, Abstände und Winkel eröffnet.

Die erste Form der Multiplikation nennt man das *Skalarprodukt* (englisch *dot product* oder *inner product*) zweier Vektoren.²² Hat man ein Koordinatensystem und dementsprechend die Komponenten der beiden Vektoren, so werden einfach alle Komponenten paarweise multipliziert und die so erhaltenen Produkte addiert. Das Ergebnis ist also ein Skalar und nicht etwa ein Vektor! Zwei Beispiele:

Skalarprodukt

$$\begin{pmatrix} 8 \\ -2 \\ 3 \end{pmatrix} \cdot \begin{pmatrix} 1 \\ 3 \\ 0 \end{pmatrix} = 8 \cdot 1 + (-2) \cdot 3 + 3 \cdot 0 = 2$$

²²Achtung: Das Skalarprodukt ist etwas ganz anderes als die [Skalarmultiplikation](#)!

$$\begin{pmatrix} 4 \\ 2 \end{pmatrix} \cdot \begin{pmatrix} -4 \\ 7 \end{pmatrix} = 4 \cdot (-4) + 2 \cdot 7 = -2$$

a · b Die übliche Schreibweise für das Skalarprodukt ist, wie wir gerade gesehen haben, ein Multiplikationspunkt, den man manchmal ganz weglässt, also **a · b** oder kurz **ab**. Dass man ein und dasselbe Zeichen für verschiedene Operationen (skalare Multiplikation, Skalarprodukt, Produkt von Skalaren, Produkt von Matrizen) verwendet, sollte Sie nicht mehr wundern. Solange für die einzelnen beteiligten Objekte klar ist, ob es sich um Skalare, Vektoren oder Matrizen handelt, kann man nicht wirklich durcheinandergeraten.²³

Lemma 28.1

Für alle Vektoren **a**, **b** und **c** aus $\mathbb{R}^n$ sowie alle Skalare $\lambda \in \mathbb{R}$ gelten die folgenden Aussagen:

- **a · b = b · a**
- **(λ · a) · b = λ · (a · b)**
- **(a + b) · c = a · c + b · c**

Aufgabe 28.1. Überzeugen Sie sich (entweder anhand von Beispielen oder abstrakt), dass die Aussagen von Lemma 28.1 korrekt sind.

a² **Aufgabe 28.2.** Berechnen Sie das Skalarprodukt von **a** = (a₁, a₂) mit sich selbst. (Dafür schreibt man manchmal auch kürzer **a²**.) Kommt Ihnen das bekannt vor?

Aufgabe 28.3. Berechnen Sie **(a – b)²** mithilfe der Regeln aus Lemma 28.1. (Sie sollen also ausmultiplizieren und zusammenfassen.)

Für den Vektor **a** = (a₁, a₂) gilt **a · a** = a₁² + a₂². Zwei Dinge fallen dabei auf:

- Der Wert kann nicht negativ sein, weil es die Summe von zwei Quadraten ist. Und null kann nur dann herauskommen, wenn **a** der Nullvektor ist. (Man sagt, dass das Skalarprodukt *positiv definit* ist.)
- Der Wert ist (nach **Pythagoras**) das Quadrat der Länge des Vektors.

Beides gilt auch dann noch, wenn wir im Raum statt in der Ebene rechnen, wobei vielleicht nicht sofort klar ist, warum für den Vektor **a** = (a₁, a₂, a₃) ∈ ℝ³ die Zahl **a · a** = a₁² + a₂² + a₃² das Quadrat seiner Länge sein soll. Machen Sie sich dafür anhand einer Skizze das folgende Vorgehen klar:

- Projizieren Sie den Vektor in die x-y-Ebene. Das ergibt den Vektor **a'** mit den Komponenten (a₁, a₂, 0). Da es sich bei **a'** faktisch um einen zweidimensionalen Vektor handelt, ist das Quadrat seiner Länge nach den bisherigen Überlegungen **a' · a'**.
- Der Vektor **a** ergibt sich als Summe aus **a'** und (0, 0, a₃) und die beiden Summanden stehen offenbar senkrecht aufeinander. Wir können also erneut

²³Manchmal kann es jedoch wirklich zu Verwechslungen kommen; wenn es nämlich noch eine *andere* Multiplikation zwischen den Vektoren gibt, für die man traditionell den üblichen Multiplikationspunkt benutzt. In diesem Fall schreibt man für das Skalarprodukt auch **⟨a, b⟩** statt **a · b**.

Pythagoras anwenden und erhalten als Quadrat der Länge von $\mathbf{a}$ den Wert $a_3^2 + \mathbf{a}' \cdot \mathbf{a}'$. Das ist $a_1^2 + a_2^2 + a_3^2$.

Wir definieren daher ganz allgemein:

Für jeden Vektor $\mathbf{a} \in \mathbb{R}^n$ definieren wir seine **Norm** durch

$$\|\mathbf{a}\| = \sqrt{\mathbf{a} \cdot \mathbf{a}}.$$

Definition

In den Räumen $\mathbb{R}^2$ und $\mathbb{R}^3$ ist die Norm eines Vektors also seine Länge.

Aufgabe 28.4. Verifizieren Sie anhand von Beispielen, dass $\|\lambda \cdot \mathbf{a}\| = |\lambda| \cdot \|\mathbf{a}\|$ für alle Vektoren $\mathbf{a}$ und alle Skalare λ gilt. Vergessen Sie dabei nicht, auch negative Werte für λ auszuprobieren.

Aufgabe 28.5. Welche Norm hat der Vektor $\mathbf{a} = (3, 4, -1)^T$? Können Sie einen Vektor angeben, der in dieselbe Richtung wie $\mathbf{a}$ zeigt, der aber die Norm 1 hat?

Aufgabe 28.6. Können Sie einen Zusammenhang zwischen Matrizenmultiplikation und Skalarprodukt erkennen? (Schauen Sie sich gegebenenfalls das [Falksche Schema](#) noch einmal an.)

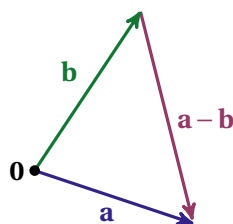
Wir fassen die wesentlichen Eigenschaften der Norm im folgenden Lemma zusammen. Die erste sollte offensichtlich sein, die zweite wurde in Aufgabe 28.4 begründet. Die dritte ist die *Dreiecksungleichung*, die Sie schon von den [komplexen Zahlen](#) kennen und die für Vektoren geometrisch auf dasselbe hinausläuft: In einem Dreieck kann die direkte Verbindung zweier Ecken durch eine Seite nicht länger sein als der „Umweg“ über die anderen beiden Seiten. (Machen Sie sich ggf. noch einmal eine Skizze.)

Für alle $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$ und alle $\lambda \in \mathbb{R}$ gilt:

- (i) $\|\mathbf{x}\| \geq 0$ und $\|\mathbf{x}\| = 0 \Leftrightarrow \mathbf{x} = \mathbf{0}$
- (ii) $\|\lambda \mathbf{x}\| = |\lambda| \cdot \|\mathbf{x}\|$
- (iii) $\|\mathbf{x} + \mathbf{y}\| \leq \|\mathbf{x}\| + \|\mathbf{y}\|$

Lemma 28.2

Wenn man die Längen von Vektoren berechnen kann, dann kann man natürlich auch Abstände bestimmen. In der folgenden Skizze ist offenbar der Abstand der Endpunkte der Vektoren $\mathbf{a}$ und $\mathbf{b}$ die Länge des Vektors $\mathbf{a} - \mathbf{b}$. (Siehe auch Seite 171 zur Subtraktion von Vektoren.)



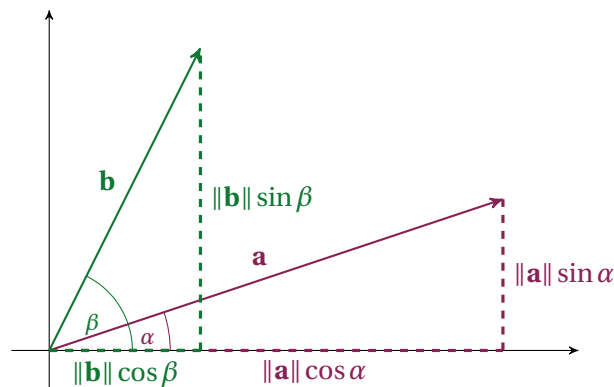
Abstand
 $d(\mathbf{p}, \mathbf{q})$

Allgemein ergibt sich also für den *Abstand* zweier Punkte $\mathbf{p}$ und $\mathbf{q}$, den wir mit $d(\mathbf{p}, \mathbf{q})$ bezeichnen werden, der Wert

$$d(\mathbf{p}, \mathbf{q}) = \|\mathbf{p} - \mathbf{q}\| = \|\mathbf{q} - \mathbf{p}\|.$$

Man nennt das manchmal auch den *euklidischen Abstand*, weil man Abstände auch anders definieren kann.

Durch den Zusammenhang zwischen dem Skalarprodukt und der Länge des Vektors haben wir bereits eine Eigenschaft des Skalarproduktes kennengelernt, die unabhängig von irgendwelchen Koordinaten ist. In diesem Abschnitt werden wir nun eine Beziehung zwischen Skalarprodukt und Winkeln entdecken. Dazu stellen wir die Komponenten von zwei Vektoren, die wir multiplizieren wollen, über den Winkel dar, den die Vektoren jeweils mit der horizontalen Achse des Koordinatensystems bilden.



Man sieht an der Skizze, dass man $\mathbf{a}$ und $\mathbf{b}$ folgendermaßen darstellen kann:

$$\mathbf{a} = \|\mathbf{a}\| \cdot \begin{pmatrix} \cos \alpha \\ \sin \alpha \end{pmatrix} \quad \mathbf{b} = \|\mathbf{b}\| \cdot \begin{pmatrix} \cos \beta \\ \sin \beta \end{pmatrix}$$

Und man erhält mithilfe der Additionstheoreme:

$$\begin{aligned} \mathbf{a} \cdot \mathbf{b} &= \|\mathbf{a}\| \cdot \|\mathbf{b}\| \cdot (\cos \alpha \cos \beta + \sin \alpha \sin \beta) \\ &= \|\mathbf{a}\| \cdot \|\mathbf{b}\| \cdot \cos(\beta - \alpha) \end{aligned} \quad (28.1)$$

Das Skalarprodukt ergibt sich also als Produkt der Längen der beiden Vektoren mit dem Kosinus des Winkels *zwischen* den beiden Vektoren. Das ist eine wichtige Aussage, weil sie zeigt, dass wir das Skalarprodukt rein geometrisch (ohne das Koordinatensystem der analytischen Geometrie) hätten definieren können. Und weil sie uns durch Umformung der obigen Gleichung eine Formel für den Winkel zwischen zwei Vektoren gibt:²⁴

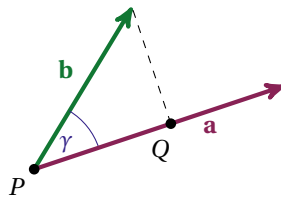
$$\angle(\mathbf{a}, \mathbf{b}) = \arccos \frac{\mathbf{a} \cdot \mathbf{b}}{\|\mathbf{a}\| \cdot \|\mathbf{b}\|}$$

²⁴Natürlich nur für den Fall, dass weder $\mathbf{a}$ noch $\mathbf{b}$ der Nullvektor ist. In diesem Fall würde es auch keinen Sinn ergeben, vom Winkel zwischen den Vektoren zu sprechen.

Genaugenommen gibt es natürlich immer *zwei* Winkel zwischen zwei Vektoren, aber da der **Arkuskosinus** nur Werte zwischen null und 180° liefert, ist mit $\angle(\mathbf{a}, \mathbf{b})$ in der obigen Formel der *kleinere* der beiden Winkel gemeint. Der andere hat offenbar den Wert $360^\circ - \angle(\mathbf{a}, \mathbf{b})$

Aufgabe 28.7. Berechnen Sie mithilfe eines Taschenrechners den Winkel zwischen den Vektoren $(32, 63)$ und $(-4, 15)$ in Grad.

Es gibt noch eine weitere Möglichkeit, das Skalarprodukt $\mathbf{a} \cdot \mathbf{b}$ geometrisch zu interpretieren. Dazu nennen wir den Winkel zwischen den beiden Vektoren γ und projizieren $\mathbf{b}$ senkrecht auf $\mathbf{a}$.



Die Länge l der Projektion (der Strecke von P nach Q) ist wegen des rechtwinkligen Dreiecks offenbar $l = \|\mathbf{b}\| \cdot \cos \gamma$. Gleichzeitig haben wir mit (28.1)

$$\|\mathbf{a}\| \cdot l = \|\mathbf{a}\| \cdot \|\mathbf{b}\| \cdot \cos \gamma = \mathbf{a} \cdot \mathbf{b}.$$

Man kann also das Skalarprodukt auch als Länge der Projektion des zweiten Vektors auf den ersten multipliziert mit der Länge des ersten interpretieren. (Wobei man beachten muss, dass anders als in der Skizze der Winkel γ auch größer als $\pi/2$ sein kann. Im Allgemeinen erhält man also die Länge der Projektion zusammen mit einem Vorzeichen, das einem sagt, ob der Winkel zwischen den Vektoren kleiner oder größer als ein rechter ist.)

Sowohl diese geometrische Interpretation als auch die vorherige funktionieren ebenso im Raum. Und beide benötigen keine Koordinaten.

Aufgabe 28.8. Gilt die obige Interpretation auch umgekehrt, wenn man also $\mathbf{a}$ auf $\mathbf{b}$ projiziert?

Die Gleichung (28.1) liefert uns noch einen weiteren Zusammenhang. Stehen zwei Vektoren senkrecht (man sagt auch: *orthogonal*) zueinander – beträgt der Winkel zwischen ihnen also 90° –, so ist der Kosinus dieses Winkels null, also muss auch ihr Skalarprodukt verschwinden. Man kann daher am Skalarprodukt erkennen, ob zwei Vektoren orthogonal zueinander sind!

orthogonal

Aufgabe 28.9. Geben Sie einen Vektor an, der orthogonal zu $\mathbf{v} = (3, 7)$ ist.

Aufgabe 28.10. Geben Sie einen Vektor an, der orthogonal zu $\mathbf{v} = (3, 7, -2)$ ist.

Aufgabe 28.11. Welchen Wert hat $(\mathbf{a} - \mathbf{b})^2$ (siehe Aufgabe 28.3), wenn $\mathbf{a}$ und $\mathbf{b}$ orthogonal zueinander sind? Welche geometrische Bedeutung hat dieses Ergebnis?

Während es (siehe Aufgabe 28.9) vergleichsweise einfach ist, einen Vektor zu finden, der senkrecht auf *einem* vorgegebenen steht, ist zunächst nicht klar, wie man im dreidimensionalen Raum einen Vektor findet, der gleichzeitig auf *zwei* gegebenen Vektoren senkrecht steht. Dafür gibt es aber eine einfache Methode, die wir uns nun ohne weitere geometrische Herleitung anschauen.

Definition

Das **Kreuz-** bzw. **Vektorprodukt** wird durch

$$\mathbf{a} \times \mathbf{b} = \begin{pmatrix} a_2 b_3 - a_3 b_2 \\ a_3 b_1 - a_1 b_3 \\ a_1 b_2 - a_2 b_1 \end{pmatrix} \quad (28.2)$$

für Vektoren des Raums $\mathbb{R}^3$ definiert.

Das so definierte Produkt ist also ein Vektor und nicht wie beim Skalarprodukt ein Skalar.

Aufgabe 28.12. Berechnen Sie $\mathbf{a} \times \mathbf{b}$ und $\mathbf{b} \times \mathbf{a}$ für $\mathbf{a} = (2, 3, -1)$ und $\mathbf{b} = (4, -2, 5)$. Was fällt Ihnen auf?

Aufgabe 28.13. Berechnen Sie das Skalarprodukt $(\mathbf{a} \times \mathbf{b}) \cdot \mathbf{a}$, wobei $\mathbf{a}$ und $\mathbf{b}$ beliebig sind, also $\mathbf{a} = (a_1, a_2, a_3)$ und $\mathbf{b} = (b_1, b_2, b_3)$ wie in der Definition (28.2).

Aufgabe 28.14. Was ergibt sich, wenn man das Vektorprodukt zweier paralleler Vektoren bildet?

Aufgabe 28.15. Sei $\mathbf{v} = (1, 2, 3)$, $\mathbf{w} = (2, a, 2)$ und $\mathbf{x} = \mathbf{v} \times \mathbf{w}$. Welchen Wert muss a haben, damit $\mathbf{x}$ senkrecht auf $\mathbf{e}_3$ steht?

Wir fassen die Eigenschaften des Vektorprodukts zusammen:

Lemma 28.3

Für alle $\mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathbb{R}^3$ und alle $\alpha \in \mathbb{R}$ gilt

- (i) $(\alpha \mathbf{x}) \times \mathbf{y} = \mathbf{x} \times (\alpha \mathbf{y}) = \alpha (\mathbf{x} \times \mathbf{y})$
- (ii) $(\mathbf{x} + \mathbf{y}) \times \mathbf{z} = (\mathbf{x} \times \mathbf{z}) + (\mathbf{y} \times \mathbf{z})$ und $\mathbf{z} \times (\mathbf{x} + \mathbf{y}) = (\mathbf{z} \times \mathbf{x}) + (\mathbf{z} \times \mathbf{y})$
- (iii) $\mathbf{x} \times \alpha \mathbf{x} = \mathbf{0}$
- (iv) $\mathbf{y} \times \mathbf{x} = -\mathbf{x} \times \mathbf{y}$

Aufgabe 28.16. Überzeugen Sie sich (durch Beispiele oder auch ganz allgemein), dass die Aussagen aus Lemma 28.3, die nicht bereits in anderen Aufgaben thematisiert wurden, richtig sind.

Rechtssystem

Die wichtigste Eigenschaft für die Praxis ist jedoch, dass $\mathbf{a} \times \mathbf{b}$ senkrecht auf $\mathbf{a}$ und $\mathbf{b}$ steht (siehe Aufgabe 28.13) und dass $\mathbf{a}$, $\mathbf{b}$ und $\mathbf{a} \times \mathbf{b}$ ein *Rechtssystem* im Sinne der *Drei-Finger-Regel* bilden.²⁵ (Letzteres bedeutet, dass eine Matrix, deren Spalten

²⁵Außerdem ist $\|\mathbf{a} \times \mathbf{b}\|$ die Fläche des von $\mathbf{a}$ und $\mathbf{b}$ aufgespannten Parallelogramms. Das werden wir jedoch im Folgenden nicht verwenden.

diese drei Vektoren in dieser Reihenfolge sind, niemals eine negative Determinante hat.)

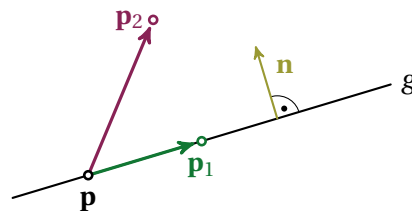
Es gibt übrigens keinen Grund, Gleichung (28.2) auswendig zu lernen. Verwenden Sie lieber die folgende Eselsbrücke. Schreiben Sie (im Geiste) die beiden zu multiplizierenden Vektoren als erste und zweite Spalte einer 3×3 -Matrix hin und schreiben Sie in die dritte Spalte die Symbole $\mathbf{e}_1$ bis $\mathbf{e}_3$. Dann berechnen Sie die Determinante dieser Matrix mit der Regel von Sarrus.

$$\begin{vmatrix} a_1 & b_1 & \mathbf{e}_1 \\ a_2 & b_2 & \mathbf{e}_2 \\ a_3 & b_3 & \mathbf{e}_3 \end{vmatrix} \\ = a_1 b_2 \mathbf{e}_3 + a_3 b_1 \mathbf{e}_2 + a_2 b_3 \mathbf{e}_1 - a_3 b_2 \mathbf{e}_1 - a_1 b_3 \mathbf{e}_2 - a_2 b_1 \mathbf{e}_3 \\ = (a_2 b_3 - a_3 b_2) \mathbf{e}_1 + (a_3 b_1 - a_1 b_3) \mathbf{e}_2 + (a_1 b_2 - a_2 b_1) \mathbf{e}_3$$

Interpretieren Sie nun $\mathbf{e}_1$ bis $\mathbf{e}_3$ als die kanonischen Einheitsvektoren, so haben Sie gerade das Kreuzprodukt berechnet.

Wir werden nun sehen, dass man Geraden in der Ebene und Ebenen im Raum auch anders als durch die Punkt-Richtungs-Form beschreiben kann. Dafür schauen wir uns in der folgenden Skizze eine Gerade g an, zu der wir einen Punkt $\mathbf{p}$ kennen, der auf g liegt, und einen Vektor $\mathbf{n}$, der senkrecht auf g steht. So einen Vektor nennt man übrigens (auch in der Computergrafik) einen *Normalenvektor* zu g (wobei „normal“ hier für rechtwinklig steht und etwas anderes bedeutet als die ähnlichen Wörter *Norm* und *normieren*).

Normalenvektor



Offensichtlich ist der Vektor von $\mathbf{p}$ zu einem Punkt $\mathbf{p}_1$, der auf g liegt, orthogonal zu $\mathbf{n}$, während der Vektor von $\mathbf{p}$ zu einem Punkt $\mathbf{p}_2$, der *nicht* auf g liegt, nicht orthogonal zu $\mathbf{n}$. Man kann also g beschreiben als die Menge aller Punkte $\mathbf{x}$, für die $\mathbf{n} \cdot (\mathbf{x} - \mathbf{p}) = 0$ gilt. Vergibt man eine Bezeichnung wie k für den Skalar $\mathbf{n} \cdot \mathbf{p}$, so erhält man eine sogenannte *Normalenform* für die Gerade:

Normalenform

$$g = \{\mathbf{x} \in \mathbb{R}^2 : \mathbf{n} \cdot \mathbf{x} - k = 0\} \quad (28.3)$$

Aufgabe 28.17. Prüfen Sie, ob die beiden Punkte $(1, 2)$ und $(-2, 4)$ auf der Geraden $h = \{(x, y) \in \mathbb{R}^2 : (3, 1) \cdot (x, y) + 2 = 0\}$ liegen.

Aufgabe 28.18. Geben Sie für $g = (-2, 1) + \mathbb{R}(2, -5)$ eine Normalenform an.

Aufgabe 28.19. Machen Sie sich, evtl. durch eine Skizze, klar, dass

$$\{\mathbf{x} \in \mathbb{R}^3 : \mathbf{n} \cdot \mathbf{x} - k = 0\} \quad (28.4)$$

eine Ebene im Raum beschreibt, wenn k eine reelle Zahl und $\mathbf{n} \neq \mathbf{0}$ ein Vektor aus der Menge $\mathbb{R}^3$ ist.

Aufgabe 28.20. Geben Sie eine Normalenform für die Ebene E an, die die drei Punkte $P_1 = (1, 2, 1)$, $P_2 = (3, -1, 5)$ und $P_3 = (4, 0, -6)$ enthält. (Hinweis: Stellen Sie zunächst eine Punkt-Richtungs-Form für E auf und verwenden Sie dann das Vektorprodukt, um einen Normalenvektor zu erhalten.)

Aufgabe 28.21. Machen Sie sich klar, dass (28.3) eine lineare Gleichung mit zwei Unbekannten ist. Ebenso ist (28.4) eine lineare Gleichung mit drei Unbekannten. Geometrisch beschreibt also ein **lineares Gleichungssystem**, das aus zwei Gleichungen mit zwei Unbekannten besteht, den Schnitt zweier Geraden in der Ebene, während eines mit drei Gleichungen und drei Unbekannten den Schnitt dreier Ebenen im Raum beschreibt. Darum gibt es im „typischen“ Fall genau eine Lösung, nämlich einen Schnittpunkt.

Aufgabe 28.22. Geben Sie einen Vektor an, der senkrecht auf den beiden Vektoren $(3, -4, 2)$ und $(-1, -1, 4)$ steht und der außerdem die Norm 3 hat.

Wenn man in der Normalenform den Normalenvektor auch noch normiert (siehe Aufgabe 28.5), so erhält man eine weitere schöne Eigenschaft. Dazu muss man nur in einer Skizze wie der auf Seite 227 sehen, dass beispielsweise die Länge der Projektion des Vektors von $\mathbf{p}$ nach $\mathbf{p}_2$ auf $\mathbf{n}$ dem Abstand des Punktes $\mathbf{p}_2$ von g entspricht. Und nach unseren Überlegungen von Seite 225 ist $\mathbf{n} \cdot (\mathbf{p}_2 - \mathbf{p})$ diese Länge, wenn $\mathbf{n}$ die Länge 1 hat. Ist $\mathbf{n}$ also normiert, so liefert der Term $\mathbf{n} \cdot \mathbf{x} - k$ aus (28.3) also den Abstand des Punktes $\mathbf{x}$ von g . (Und ebenso gilt das für (28.4) und den Abstand zur dargestellten Ebene.) Man kann in diesem Fall die Normalenform also auch interpretieren als Beschreibung der Menge der Punkte, deren Abstand von der Geraden bzw. Ebene null ist.

Aufgabe 28.23. Bei der Erklärung im obigen Absatz wurde ein Detail unterschlagen. Welches?

Aufgabe 28.24. Geben Sie für die Gerade g aus Aufgabe 28.18 eine Normalenform mit normiertem Normalenvektor an.

Aufgabe 28.25. Berechnen Sie den Abstand des Punktes $(4, -3)$ von der Geraden g aus Aufgabe 28.24.

Aufgabe 28.26. Wir wissen bereits, dass es viele verschiedene Normalenformen derselben Geraden bzw. Ebene gibt. Wenn man fordert, dass der Normalenvektor normiert ist, gibt es aber immer nur eine, oder?

hessesche
Normalenform

Die eindeutig bestimmte Normalenform, die man erhält, wenn in (28.3) erstens $\mathbf{n}$ normiert und zweitens k nicht negativ ist, nennt man die *hessesche Normalenform*. In dieser Form liefert (der Betrag von) $\mathbf{n} \cdot \mathbf{x} - k$ nicht nur den Abstand von $\mathbf{x}$ zur Geraden (bzw. Ebene), sondern man kann am Vorzeichen dieses Terms auch noch erkennen, ob $\mathbf{x}$ auf derselben Seite der Geraden (oder Ebene) wie $\mathbf{0}$ liegt. Das ist nämlich genau dann der Fall, wenn der Term negativ ist.

Aufgabe 28.27. Ist die Lösung von Aufgabe 28.24 in hessescher Normalenform angegeben?

Aufgabe 28.28. Geben Sie für die Ebene E durch die $P_1 = (-3, -2, 4)$, $P_2 = (3, 1, 2)$ und $P_3 = (2, 2, 3)$ die hessesche Normalenform an. Von den drei Punkten $Q_1 = (2, -2, 2)$, $Q_2 = (-2, 2, 2)$ und $Q_3 = (3, -2, 1)$ liegt einer nicht auf derselben Seite von E wie die beiden anderen. Geben Sie den Abstand dieses Punktes von E an.

Bemerkungen

Beachten Sie, dass das Vektorprodukt, im Gegensatz zu all den anderen Dingen, die wir bisher mit Vektoren angestellt haben, *ausschließlich* für Vektoren im Raum $\mathbb{R}^3$ definiert ist. (Es gibt eine Verallgemeinerung für $\mathbb{R}^n$ für beliebige n , aber die hier gezeigte Version für $n = 3$ ist die mit Abstand am häufigsten verwendete.)

Da ich das schon öfter gelesen habe: Es heißt *hessesche* und nicht *hessische* Normalenform, weil diese Form nach dem Mathematiker [Otto Hesse](#) und nicht nach dem Bundesland Hessen benannt ist.

Warum *definiert* man die Norm, wenn damit doch einfach die Länge gemeint ist? Könnte man es nicht einfach *Länge* nennen? Die Antwort ist, dass man den Begriff des Skalarproduktes verallgemeinern kann. Man nennt in der Mathematik jede Verknüpfung *ein* Skalarprodukt, die a) einem Paar *abstrakter* Vektoren²⁶ einen Skalar zuordnet, b) die in Lemma 28.1 aufgezählten Eigenschaften hat und c) positiv definit ist. Und wenn man ein Skalarprodukt hat, kann man immer auch eine zugehörige Norm definieren. Das ist dann aber zunächst nur ein abstrakter Wert, den man zwar für Analogieschlüsse wie eine Länge behandeln kann, der aber mit einer Länge im geometrischen Sinne nichts mehr zu tun haben muss. Genauso macht man es auch mit den Abständen und den Winkeln. Man geht immer von der (euklidischen) Geometrie in den uns bekannten anschaulichen Räumen aus und abstrahiert dann.

★ Skalar- und Vektorprodukt im Computer

Hier ein paar einfache Definition für die Operationen aus diesem Kapitel. Beachten Sie, dass wir das Skalarprodukt als dotProd [bereits definiert](#) hatten.

```
from math import sqrt

norm = lambda X: sqrt(dotProd(X, X))
# Normalisieren eines Vektors
normalize = lambda V: scalarVectorMult(1/norm(V), V)
# Hilfsfunktion für die Subtraktion von Vektoren
vectorSub = lambda A, B: [a - b for a, b in zip(A, B)]
# Abstand
```

²⁶Das können Funktionen oder andere Objekte sein. Man muss mit ihnen nur wie mit Vektoren rechnen können. Im schon erwähnten Video [Was ist lineare Algebra?](#) wird dieser Abstraktionsvorgang erläutert.

```

distance = lambda A,B: norm(vectorSub(A,B))
# Vektorprodukt
def crossProduct(A,B):
    return [A[1] * B[2] - A[2] * B[1],
            A[2] * B[0] - A[0] * B[2],
            A[0] * B[1] - A[1] * B[0]]

```

29. Raumdrehungen und Quaternionen

Weil Drehungen im Raum in der Computergrafik wichtig sind, gehen wir in diesem Kapitel kurz auf sie ein. Insbesondere schauen wir uns in diesem Zusammenhang auch die sogenannten *Quaternionen* an, die in vielen Programmen zur Darstellung von Raumdrehungen verwendet werden. Diese sogenannten "[hyperkomplexen Zahlen](#)" wurden explizit im 19. Jahrhundert „erfunden“, um solche Drehungen darzustellen. Wir werden sie relativ knapp behandeln, aber bei Interesse gibt es ausführlichere Informationen zu diesem Thema in den Videos [Was sind und was sollen die Quaternionen?](#) sowie [Was Sie schon immer über Drehungen wissen wollten](#).

Definition

Eine lineare Abbildung f von $\mathbb{R}^n$ nach $\mathbb{R}^n$ heißt **orthogonal**, wenn für alle $\mathbf{v}, \mathbf{w} \in \mathbb{R}^n$ die Beziehung $f(\mathbf{v}) \cdot f(\mathbf{w}) = \mathbf{v} \cdot \mathbf{w}$ gilt. Und eine $n \times n$ -Matrix wird **orthogonal** genannt, wenn sie eine orthogonale lineare Abbildung darstellt.

Orthogonale Abbildungen erhalten also das Skalarprodukt. Damit erhalten sie gemäß [dem letzten Abschnitt](#) auch Normen, Abstände und Winkel:

$$\begin{aligned}
 \|f(\mathbf{v})\| &= \sqrt{f(\mathbf{v}) \cdot f(\mathbf{v})} = \sqrt{\mathbf{v} \cdot \mathbf{v}} = \|\mathbf{v}\| \\
 \|f(\mathbf{v}) - f(\mathbf{w})\| &= \|f(\mathbf{v} - \mathbf{w})\| = \|\mathbf{v} - \mathbf{w}\| \\
 \angle(f(\mathbf{v}), f(\mathbf{w})) &= \arccos \frac{f(\mathbf{v}) \cdot f(\mathbf{w})}{\|f(\mathbf{v})\| \|f(\mathbf{w})\|} = \arccos \frac{\mathbf{v} \cdot \mathbf{w}}{\|\mathbf{v}\| \|\mathbf{w}\|} = \angle(\mathbf{v}, \mathbf{w})
 \end{aligned}$$

Aufgabe 29.1. Wenn $\mathbf{A}$ eine orthogonale Matrix ist, was kann man dann über ihre [Determinante](#) sagen? (Hinweis: Erinnern Sie sich, dass Determinanten Flächen- bzw. Volumenfaktoren sind.)

Stellt man wie in Aufgabe [28.6](#) rein formal das Skalarprodukt als Produkt von Matrizen dar, so erhält man für $\mathbf{A} \in \mathbb{R}^{n \times n}$ und $\mathbf{v}, \mathbf{w} \in \mathbb{R}^n$ mit [\(21.2\)](#)

$$(\mathbf{A}\mathbf{v}) \cdot (\mathbf{A}\mathbf{w}) = (\mathbf{A}\mathbf{v})^T \cdot (\mathbf{A}\mathbf{w}) = \mathbf{v}^T \cdot (\mathbf{A}^T \cdot \mathbf{A}) \cdot \mathbf{w}, \quad (29.1)$$

wobei der linke Multiplikationspunkt für das Skalarprodukt steht und alle anderen die Matrizenmultiplikation repräsentieren sollen. Soll der Term in [\(29.1\)](#) für *alle* $\mathbf{v}$ und $\mathbf{w}$ identisch mit dem Skalarprodukt $\mathbf{v} \cdot \mathbf{w}$ bzw. dem Matrizenprodukt $\mathbf{v}^T \cdot \mathbf{w}$ sein, so ist das nur möglich, wenn das Produkt in Klammern die Einheitsmatrix $\mathbf{E}_n$ ist:

Eine Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ ist genau dann orthogonal, wenn $\mathbf{A}^T \cdot \mathbf{A} = \mathbf{A} \cdot \mathbf{A}^T = \mathbf{E}_n$ gilt.

Lemma 29.1

Das hat zwei wichtige Konsequenzen. Um zu prüfen, ob eine Matrix orthogonal ist, muss man sie nur mit der transponierten Matrix multiplizieren. Und wenn man von einer Matrix bereits weiß, dass sie orthogonal ist, dann muss man für das Berechnen der Inversen nicht mühsam wie in Abschnitt 26 vorgehen, sondern kann einfach nur transponieren.

Aufgabe 29.2. Überprüfen Sie für

$$\mathbf{A} = \frac{1}{2} \cdot \begin{pmatrix} 1 & \sqrt{3} \\ \sqrt{3} & 1 \end{pmatrix} \quad \text{und} \quad \mathbf{B} = \frac{1}{15} \cdot \begin{pmatrix} -11 & 10 & -2 \\ 10 & 10 & -5 \\ -2 & -5 & -14 \end{pmatrix},$$

ob die Matrizen orthogonal sind.

Aufgabe 29.3. Geben Sie eine 2×2 -Matrix an, die die Determinante 1 hat, die aber *nicht* orthogonal ist.

Eine Funktion²⁷ f von $\mathbb{R}^n$ nach $\mathbb{R}^n$ wird **Isometrie** genannt, wenn sie Abstände nicht ändert, wenn also

$$d(\mathbf{v}, \mathbf{w}) = \|\mathbf{v} - \mathbf{w}\| = \|f(\mathbf{v}) - f(\mathbf{w})\| = d(f(\mathbf{v}), f(\mathbf{w}))$$

für alle $\mathbf{v}, \mathbf{w} \in \mathbb{R}^n$ gilt.

Definition

Einfache Beispiele für Isometrien sind **Translationen**: Betrachten wir die Verschiebung τ um den Vektor $\mathbf{a}$, also $\tau(\mathbf{v}) = \mathbf{v} + \mathbf{a}$, dann gilt natürlich

$$d(\tau(\mathbf{v}), \tau(\mathbf{w})) = \|(\mathbf{v} + \mathbf{a}) - (\mathbf{w} + \mathbf{a})\| = \|\mathbf{v} - \mathbf{w}\| = d(\mathbf{v}, \mathbf{w}).$$

Anschaulich ist zudem klar, dass Drehungen und Spiegelungen Isometrien sind. Es wird sich herausstellen, dass damit schon alle Isometrien aufgezählt sind.

Die Beweise für die drei folgenden Aussagen findet man im alten Skript, das am Ende der „**Gebrauchsanweisung**“ erwähnt wird.

Eine lineare Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^n$ ist genau dann eine Isometrie, wenn sie orthogonal ist.

Lemma 29.2

Jede Isometrie von $\mathbb{R}^n$ nach $\mathbb{R}^n$ lässt sich als Komposition einer orthogonalen linearen Abbildung und einer Translation darstellen. Bildet insbesondere so eine Isometrie $\mathbf{0}$ auf $\mathbf{0}$ ab, dann ist sie eine orthogonale lineare Abbildung.

Korollar 29.3

²⁷Es wird nicht gefordert, dass f eine lineare Abbildung ist!

Korollar 29.3 ist gewissermaßen eine Rechtfertigung dafür, warum es für viele Anwendungen in der Computergrafik ausreicht, sich wie in Abschnitt 25 auf Translationen und einige lineare Abbildungen zu beschränken.

Satz 29.4

Die orthogonalen Abbildungen von $\mathbb{R}^2$ nach $\mathbb{R}^2$ sind genau die Drehungen und die Spiegelungen.

Drehspiegelung

Dieser Satz lässt sich verallgemeinern: Unabhängig von der Dimension gibt es immer nur zwei Typen von orthogonalen Abbildungen: die Drehungen und die sogenannten *Drehspiegelungen*, die sich als Kombination aus Drehung und Spiegelung darstellen lassen. Drehungen haben die Determinante 1, bei Drehspiegelungen ist die Determinante -1 . In der Ebene sind die Drehspiegelungen so simpel, dass sie einfach Spiegelungen sind, in höheren Dimensionen ist es komplizierter.

Aufgabe 29.4. Schreiben Sie eine allgemeine Drehmatrix für die Ebene hin und verifizieren Sie, dass diese Matrix orthogonal ist.

Aufgabe 29.5. $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ sei die Drehung um den Nullpunkt um den Winkel $\pi/4$ und $\mathbf{a} = (3, 4)$. Geben Sie $\|f(\mathbf{a})\|$ an.

Aufgabe 29.6. Was kann man für a einsetzen, damit diese Matrix orthogonal ist?

$$\frac{1}{\sqrt{2}} \cdot \begin{pmatrix} a & 0 & -1 \\ -1 & 0 & -1 \\ 0 & \sqrt{2} & 0 \end{pmatrix}$$

symmetrisch

Aufgabe 29.7. Eine Matrix $\mathbf{M}$ heißt *symmetrisch*, wenn $\mathbf{M} = \mathbf{M}^T$ gilt. Finden Sie durch Probieren eine 2×2 -Matrix, die sowohl orthogonal als auch symmetrisch ist und die *nicht* die Einheitsmatrix ist.

Aufgabe 29.8. $\mathbf{A}$ sei eine Drehmatrix. $\mathbf{B}$ und $\mathbf{C}$ seien Matrizen mit den Determinanten 6 und 24. Geben Sie die Determinante von $\mathbf{A}^{-1} \cdot \mathbf{B}^{-1} \cdot \mathbf{C}^T$ an. (Alle drei Matrizen sollen natürlich quadratisch und vom selben Typ sein.)

Definition

Ein Objekt der Form $u = a + bi + cj + dk$ mit reellen Zahlen a bis d nennt man ein (oder auch eine) **Quaternion**.²⁸ Dabei soll i die *imaginäre Einheit* der komplexen Zahlen sein, während j und k zwei *andere* Quaternionen sind, die weder reell noch komplex sind. Die reelle Zahl $\operatorname{Re}(u) = a$ nennt man den **Realteil** von u . Verschwindet dieser, so nennt man u **rein imaginär** oder **vektoriell**. Für die Menge aller Quaternionen schreiben wir $\mathbb{H}$.

Da das Symbol $\mathbb{Q}$ bereits vergeben ist, verwendet man $\mathbb{H}$ zu Ehren des irischen Mathematikers **William Rowan Hamilton**, der die Quaternionen zwar nicht als erster Mensch beschrieben, aber wesentlich zu ihrer Verbreitung beigetragen hat. Mit $c = d = 0$ sieht man, dass $\mathbb{C} \subseteq \mathbb{H}$ gilt.

Das Rechnen mit Quaternionen definieren wir ähnlich wie das mit komplexen Zahlen:

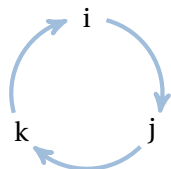
Quaternionen werden beim Rechnen mit den vier Grundrechenarten so behandelt, als wären i , j und k reelle Zahlen. Dabei wird jedoch die Regel

$$i^2 = j^2 = k^2 = ijk = -1 \quad (29.2)$$

beachtet und bei Multiplikationen darf die Reihenfolge der Symbole i , j und k nicht verändert werden.

Definition

Aus (29.2) folgt, dass $-i$, $-j$ und $-k$ die Kehrwerte von i , j und k sind. Multipliziert man die Gleichung $ijk = -1$ von rechts mit $-k$, so erhält man $ij = k$. Mit ähnlichen Umformungen kann man sich für alle Produkte der drei Symbole i , j und k herleiten, was herauskommen muss. Das kann man sich mit der folgenden Eselsbrücke merken:



Das Produkt zweier dieser Zahlen ist immer die dritte. Multipliziert man gegen die Pfeilrichtung, so kommt noch ein Minus vor das Ergebnis, also z. B. $kj = -i$.

Wir berechnen exemplarisch das Produkt von $u = 2 + 3i - 5j + k$ und $v = -4 - i - 4j + 2k$. Wenn man das von Hand macht, ist eine Tabelle sinnvoll, weil man 16 Teilprodukte bilden muss:

	-4	$-i$	$-4j$	$2k$
2	-8	$-2i$	$-8j$	$4k$
$3i$	$-12i$	3	$-12k$	$-6j$
$-5j$	$20j$	$-5k$	-20	$-10i$
k	$-4k$	$-j$	$4i$	-2

Nun werden jeweils vier „gleichartige“ Terme zusammengefasst. Beispielsweise haben wir $-2i$, $-12i$, $4i$ und $-10i$, was zusammen $-20i$ ergibt. Insgesamt erhält man $uv = -27 - 20i + 5j - 17k$.

Aufgabe 29.9. Berechnen Sie vu für die Zahlen aus dem Beispiel.

Aufgabe 29.10. Üben Sie das Multiplizieren von Quaternionen. In [Wolfram Alpha](#) kann man z. B. zum Überprüfen seiner Ergebnisse so etwas wie

```
quaternion(2,3,-5,1)*quaternion(-4,-1,-4,2)
```

eingeben.

Die folgenden Definitionen und Eigenschaften ähneln stark denen der komplexen Zahlen.

Definition

Der **Absolutbetrag** bzw. **Betrag** eines Quaternions $u = a + bi + cj + dk$ ist die reelle Zahl

$$|u| = \sqrt{a^2 + b^2 + c^2 + d^2}.$$

Gilt $|u| = 1$, so nennt man u ein **Einheitsquaternion**. Und das Quaternion $\bar{u} = a - bi - cj - dk$ wird als das zu u **konjugierte** Quaternion bezeichnet.

Lemma 29.5

Für alle $u, v \in \mathbb{H}$ gilt:

- (i) $|u| \geq 0$
- (ii) $|u| = 0 \Leftrightarrow u = 0$
- (iii) $|uv| = |u| \cdot |v|$
- (iv) $|u + v| \leq |u| + |v|$

Punkt (iv) ist mal wieder die Dreiecksungleichung und drei der vier Punkte folgen direkt aus Lemma 28.2, weil der Betrag von Quaternionen wie die Norm von $\mathbb{R}^4$ -Vektoren definiert ist. Lediglich Punkt (iii) müsste explizit nachgerechnet werden. Er zeigt, dass der Betrag sich wie in $\mathbb{C}$ und $\mathbb{R}$ mit der Multiplikation „verträgt“.

Wie das folgende Lemma zeigt, gibt es auch *fast* eine Division. Aber nicht ganz! Siehe dazu Aufgabe 29.13.

Lemma 29.6

Für alle $u \in \mathbb{H}$ gilt $u\bar{u} = \bar{u}u = |u|^2$. Damit ist $u^{-1} = 1/|u|^2 \cdot \bar{u}$ für $u \in \mathbb{H} \setminus \{0\}$ der **Kehrwert** von u , d. h., es gilt $u^{-1}u = uu^{-1} = 1$. Insbesondere ist der Kehrwert von Einheitsquaternionen besonders einfach zu berechnen: Aus $|u| = 1$ folgt $u^{-1} = \bar{u}$.

Aufgabe 29.11. Sei $u = 3 - 2i + 5j - k$. Berechnen Sie $\bar{u}$, $|u|$ und u^{-1} .

Aufgabe 29.12. Überprüfen Sie, dass die Formel $u\bar{u} = |u|^2$ korrekt ist.

Aufgabe 29.13. Sei u die Zahl aus Aufgabe 29.11 und $v = 1 + 2j + k$. Berechnen Sie $u^{-1}v$ und vu^{-1} .

Die *vektoriellen* Quaternionen werden so genannt, weil man sie häufig mit Vektoren identifiziert. Damit ist gemeint, dass man das Quaternion $u = u_1i + u_2j + u_3k$ je nach Kontext als den Vektor $\mathbf{u} = (u_1, u_2, u_3)$ betrachtet und umgekehrt. Mit einem zweiten vektoriellen Quaternion $v = v_1i + v_2j + v_3k$ ergibt sich nun (rechnen Sie nach!)

$$\begin{aligned} uv &= -(u_1v_1 - u_2v_2 - u_3v_3) \\ &\quad + (u_2v_3 - u_3v_2)i + (u_3v_1 - u_1v_3)j + (u_1v_2 + u_2v_1)k \end{aligned} \quad (29.3)$$

und das kann man in diesem Sinne interpretieren²⁹ als die Summe aus der reellen Zahl $-\mathbf{u} \cdot \mathbf{v}$ (wobei das Skalarprodukt gemeint ist) und dem Vektorprodukt $\mathbf{u} \times \mathbf{v}$.

Zwei offensichtliche Folgerungen aus (29.3) sind, dass erstens $uv = \mathbf{u} \times \mathbf{v}$ gilt, wenn u und v als Vektoren orthogonal sind, und dass zweitens $u^2 = -1$ gilt, wenn u ein Einheitsquaternion ist.

Sei nun $t = t_0 + t_1\mathbf{i} + t_2\mathbf{j} + t_3\mathbf{k}$ ein Einheitsquaternion und $\mathbf{t}'$ der Vektor (t_1, t_2, t_3) . Dann haben wir

$$1 = |t|^2 = t_0^2 + (t_1^2 + t_2^2 + t_3^2) = t_0^2 + \|\mathbf{t}'\|^2$$

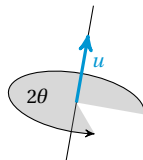
und damit ergibt sich nach dem [umgekehrten Satz des Pythagoras](#) ein rechtwinkliges Dreieck mit den Katheten t_0 und $\|\mathbf{t}'\|$. Da die Hypotenuse die Länge 1 hat, kann man dieses Dreieck in den Einheitskreis legen und erhält wie in Abschnitt 18 einen Winkel θ mit $t_0 = \cos \theta$ und $\|\mathbf{t}'\| = \sin \theta$. Durch Normieren des vektoriellen Teils von t erhält man ein vektorielles Einheitsquaternion u , so dass insgesamt

$$t = \cos \theta + u \sin \theta$$

gilt. Man kann nun zeigen, dass die durch

$$q \mapsto f_t(q) = tq\bar{t} = tq\bar{t}^{-1}$$

definierte Abbildung f_t von $\mathbb{H}$ nach $\mathbb{H}$ eine Drehung um die Achse $\mathbb{R}\mathbf{u}$ und den Winkel 2θ beschreibt, wenn q vektoriell ist und als Vektor interpretiert wird. (Man betrachtet f_t also quasi als Abbildung von $\mathbb{R}^3$ nach $\mathbb{R}^3$.)



Die Begründung dafür wird in dem [am Anfang des Abschnitts](#) erwähnten Video über Quaternionen vorgerechnet. Wir halten für die Praxis fest:

Raumdrehungen durch Quaternionen

Ist $\mathbf{u} = (u_1, u_2, u_3) \neq \mathbf{0}$ ein $\mathbb{R}^3$ -Vektor, θ ein Winkel und t das Quaternion

$$t = \cos(\theta/2) + \sin(\theta/2) \cdot \|\mathbf{u}\|^{-1} \cdot (u_1\mathbf{i} + u_2\mathbf{j} + u_3\mathbf{k}),$$

so wird durch $q \mapsto f_t(q) = tq\bar{t}$ eine Funktion definiert, die eine Raumdrehung um den Winkel θ um die Achse $\mathbb{R}\mathbf{u}$ beschreibt, wenn q vektoriell ist und q und $f_t(q)$ als Vektoren interpretiert werden.

Der Drehung in Gegenrichtung entspricht $q \mapsto \bar{t}qt$. Und beschreibt $s \in \mathbb{H}$ ebenfalls eine Drehung, so ergibt sich die Komposition $f_s \circ f_t$ als $q \mapsto (st)q(\overline{st})$; sie kann also durch ein Quaternion (nämlich st) dargestellt werden.

Satz 29.7

²⁹Tatsächlich sind historisch Skalarprodukt und Vektorprodukt so entstanden – durch Aufteilen des Produktes vektorieller Quaternionen.

Aufgabe 29.14. Berechnen Sie gemäß Satz 29.7 das Quaternion t , das die Drehung um $\mathbb{R}(3, 0, -4)$ um den Winkel $\pi/2$ beschreibt. Berechnen Sie dann, wohin der Vektor $(2, -1, 5)$ gedreht werden würde.

★**Aufgabe 29.15.** In dem [am Anfang des Abschnitts](#) erwähnten Video über Raumdrehungen erfährt man, wie eine allgemeine Drehmatrix für den Raum $\mathbb{R}^3$ aussieht, und man kann es auch [auf Wikipedia](#) nachschlagen. Sie können so eine Matrix mittels Satz 29.7 aber auch selbst berechnen, indem Sie die Vektoren $\mathbf{e}_1$ bis $\mathbf{e}_3$ drehen, um die Spalten der Matrix zu erhalten. Vielleicht sollten Sie das zur Übung zumindest für ein paar Einträge der Matrix mal machen. Der Einfachheit halber können Sie dabei annehmen, dass $\mathbf{u}$ bereits normiert ist.

Bemerkungen

Beachten Sie, dass die Multiplikation von Quaternionen im Allgemeinen nicht kommutativ ist und dass es für Quaternionen keine vollwertige Division wie in $\mathbb{C}$ oder $\mathbb{R}$ gibt. (Siehe dazu Aufgabe 29.13.)

★ Quaternionen im Computer

Wenn man ein Quaternion $a + bi + cj + dk$ in PYTHON als Liste $[a, b, c, d]$ darstellt, kann man einige der [weiter oben](#) bereits definierten Funktionen verwenden und benötigt nur noch zwei neue:

```
# Konjugieren eines Quaternions
conj = lambda q: [q[0]]+[-q[i] for i in range(1,4)]
# Multiplikation zweier Quaternionen
def mult(Q,R):
    P = [0,0,0,0]
    for i in range(4):
        for j in range(4):
            P[abs(i-j) if i*j==0 or i==j else 6-i-j] \
                += Q[i]*R[j]*(1 if i*j==0 or i%3+1==j else -1)
    return P
```

Die Rotation im Raum ist nun ganz einfach:

```
from math import cos, sin, pi

def rotate(P,u,theta):
    Q = [cos(theta/2)]+scalarVectorMult(sin(theta/2),normalize(u))
    return mult(mult(Q,[0]+P),conj(Q))[1:]
```

Zum Beispiel `rotate([4,2,0],[0,0,-1],-pi/2)` sollte (bis auf Rundungsfehler) den Vektor $(-2, 4, 0)$ liefern.

VII

Analysis

Analysis ist das Teilgebiet der Mathematik, das u. a. die Differential- und Integralrechnung umfasst, die man bereits in der Schule kennenlernt. Im Vergleich zu den anderen beiden großen Bereichen – **Geometrie** und Algebra – ist Analysis das jüngste Mitglied, das es in der heutigen Form erst seit dem Ende des 17. Jahrhunderts gibt. In den folgenden Abschnitten wiederholen und vertiefen wir einerseits Schulwissen, erarbeiten aber auch Konzepte, die in der Schule in der Regel nicht gelehrt werden.

30. Folgen

In der Analysis geht es eigentlich immer in der einen oder anderen Form um **Funktionen**, deren Funktionswerte reelle oder komplexe Zahlen sind (oder später auch **Tupel** solcher Zahlen). In diesem Abschnitt beginnen wir mit einer speziellen Sorte solcher Funktionen.

Eine **Folge** ist eine **Funktion** f , deren **Definitionsbereich** $\text{dom}(f)$ die Menge $\mathbb{N}$ ist. Man kann sich vorstellen, dass eine Folge eine „Richtung“ hat und die Funktionswerte, die man die **Folgenglieder** nennt, der Reihe nach aufgezählt werden: $f(0)$, $f(1)$, $f(2)$ und so weiter. Ist der Wertebereich einer Folge eine Teilmenge von $\mathbb{R}$ bzw. $\mathbb{C}$, so spricht man auch von einer **reellen** bzw. einer **komplexen Folge**.

Für Folgen verwendet man die abkürzende Schreibweise $(f(n))_{n \in \mathbb{N}}$ oder manchmal sogar nur $(f(n))$, wobei die zweite nur verwendet werden sollte, wenn es keine Gefahr von Missverständnissen gibt. Häufig nutzt man auch indizierte Buchstaben wie z. B. $(a_n)_{n \in \mathbb{N}}$, wobei a_n dann für den Funktionswert zum Argument n steht.

Definition

Die Funktionen

$$: \left\{ \begin{array}{l} \mathbb{N} \rightarrow \mathbb{Q} \\ n \mapsto \frac{n}{n+1} \end{array} \right. \quad \text{und} \quad : \left\{ \begin{array}{l} \mathbb{N} \rightarrow \mathbb{C} \\ n \mapsto i^n \end{array} \right.$$

sind beispielsweise Folgen, die man in obiger Schreibweise

$$\left(\frac{n}{n+1}\right)_{n \in \mathbb{N}} \quad \text{bzw.} \quad (i^n)_{n \in \mathbb{N}}$$

schreiben würde. Die Folgenglieder der linken Folge sind rationale Zahlen; die ersten vier sind 0, 1/2, 2/3 und 3/4. Die rechte Folge hat komplexe Folgenglieder. Sie beginnt mit 1, i, -1 und -i und ab dann wiederholen sich diese vier Werte zyklisch.

Aufgabe 30.1. Sei p_i die i -te Primzahl, also $p_1 = 2$, $p_2 = 3$ und so weiter. Geben Sie die ersten fünf Folgenglieder der Folge $(2^{p_{n+1}-1})_{n \in \mathbb{N}}$ an.

Aufgabe 30.2. Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge, deren erste Folgenglieder $a_0 = 0$, $a_1 = 2/3$, $a_2 = 4/5$, $a_3 = 6/7$ und $a_4 = 8/9$ sind. Geben Sie eine Rechenvorschrift für a_n an, die zu diesem Anfang passt.

Bei Folgen wird uns in der Regel eher interessieren, was „am Ende“ passiert. Daher lassen wir, um uns Schreibarbeit zu sparen, auch formal nicht ganz korrekte Folgen wie $(1/n)_{n \in \mathbb{N}}$ zu, bei denen eventuell ein paar Werte am Anfang nicht definiert sind.¹

Da Folgen immer unendlich viele Folgenglieder haben, ergibt die folgende Definition Sinn.

Definition

Von einer Aussageform $A(x)$ sagen wir, dass sie für **fast alle** Glieder einer Folge $(a_n)_{n \in \mathbb{N}}$ wahr ist, wenn es nur endlich viele „Ausnahmen“ gibt, wenn also die Menge $\{n \in \mathbb{N} : \neg A(a_n)\}$ endlich ist.

Wir würden also beispielsweise sagen, dass fast alle Folgenglieder von $(10^{9-n})_{n \in \mathbb{N}}$ kleiner als 1/1000 und fast alle Folgenglieder von $(2^n)_{n \in \mathbb{N}}$ durch 4 teilbar sind. (Die Aussageformen sind hier „ $x < 1/1000$ “ und „ $4 \mid x$ “.)

Aufgabe 30.3. Sind fast alle Folgenglieder von $(2n)_{n \in \mathbb{N}}$ gerade?

Aufgabe 30.4. Sind fast alle Folgenglieder von $(n)_{n \in \mathbb{N}}$ gerade?

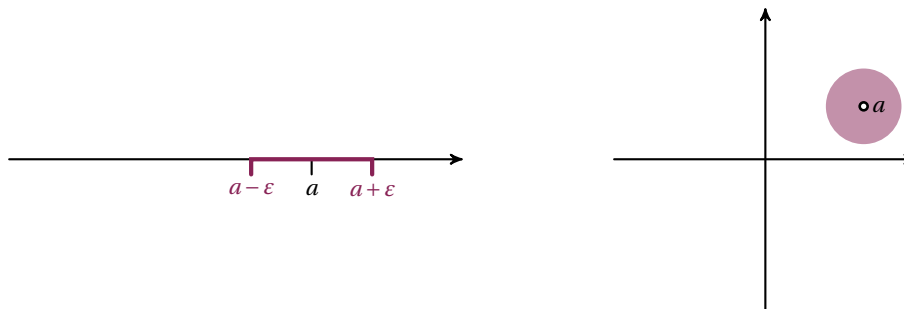
Aufgabe 30.5. Machen Sie sich klar, dass $A(x)$ genau dann für fast alle Glieder von $(a_n)_{n \in \mathbb{N}}$ gilt, wenn es eine Zahl $n_0 \in \mathbb{N}$ gibt, so dass $A(a_n)$ für $n \geq n_0$ immer wahr ist.

Definition

Eine komplexe² Folge $(a_n)_{n \in \mathbb{N}}$ **konvergiert** (oder „geht“) gegen eine Zahl $a \in \mathbb{C}$, wenn für jedes $\varepsilon > 0$ für fast alle Folgenglieder $|a_n - a| < \varepsilon$ gilt. Man schreibt dann $a_n \xrightarrow{n \in \mathbb{N}} a$ oder manchmal noch kürzer $a_n \rightarrow a$. Wenn eine Folge konvergiert, bezeichnet man sie als **konvergent**, anderenfalls als **divergent**.

¹Im konkreten Fall geht es um den Wert für $n = 0$, weil durch null dividiert werden müsste. Mit „ein paar“ Werten sind endlich viele gemeint.

Für reelle Folgen bedeutet das, dass wie im linken Teil der folgenden Skizze fast alle Folgenglieder im roten Intervall $(a - \varepsilon, a + \varepsilon)$ liegen – und zwar für *jedes* noch so kleine ε , solange es sich bei ε um eine positive Zahl handelt. (Natürlich werden es evtl. „weniger“ Glieder sein, wenn das Intervall kleiner wird, aber es werden immer *fast alle* sein.)



Für komplexe Folgen bedeutet die Definition, dass sich (im rechten Teil der Skizze) in dem roten Kreis mit Mittelpunkt a und Radius ε fast alle Folgenglieder befinden – auch hier wieder unabhängig davon, wie klein der Kreis wird. (Wenn Sie übrigens so etwas wie „für jedes $\varepsilon > 0$ “ lesen, dann kann mit ε nur eine *reelle* Zahl gemeint sein, weil die Aussage für komplexe Zahlen keinen Sinn ergeben würde.)

Als erstes Beispiel schauen wir uns vier Folgen an, die alle gegen 0 konvergieren:

- Die Folge $(1/n)_{n \in \mathbb{N}}$ konvergiert gegen 0, weil der Abstand $|1/n - 0| = 1/n$ beliebig klein wird, wenn n groß genug ist. Konkret: Ist eine Zahl $\varepsilon > 0$ vorgegeben, dann setzt man $n_0 = \lceil 1/\varepsilon \rceil$. Dann gilt $1/n_0 < \varepsilon$ und „erst recht“ $1/n < \varepsilon$ für $n > n_0$. Es gibt also nur n_0 „Ausnahmen“, für die $|1/n - 0| < \varepsilon$ *nicht* gilt.
- Die Folge $((-1)^n/n)_{n \in \mathbb{N}}$ konvergiert ebenfalls gegen 0, weil der Abstand $|(-1)^n/n - 0|$ denselben Wert wie bei der vorherigen Folge hat. Der Unterschied ist lediglich, dass die erste Folge sich auf der Zahlengeraden der Null von rechts nähert, während die Werte der zweiten Folge abwechselnd rechts und links von der Null liegen. (Der Term $(-1)^n$ sorgt offensichtlich für das wechselnde Vorzeichen. Man spricht in so einem Fall von einer *alternierenden* Folge.)
- Die Folge $(i^n/n)_{n \in \mathbb{N}}$ konvergiert gegen 0, weil der Abstand $|i^n/n - 0|$ ebenfalls denselben Wert wie bei den beiden vorherigen Folgen hat. Machen Sie sich durch eine Skizze und das Berechnen der ersten Folgenglieder klar, dass die Werte gewissermaßen eine enger werdende Spirale um die Null bilden.
- Die Folge $(0)_{n \in \mathbb{N}}$ (also die konstante Funktion $n \mapsto 0$) konvergiert gegen 0. Das ist trivial, weil der Abstand, um den es geht, einfach $|0 - 0| = 0$ ist und dieser Abstand natürlich kleiner als jedes positive ε ist. Es liegen also sogar *alle* Folgenglieder „dicht genug“ bei der Null und nicht nur fast alle. (Interessanterweise gibt es in jedem Semester Studierende, die meinen, dass diese Folge nicht gegen 0 konvergiert. Falls Sie zu denen gehören, dann schauen

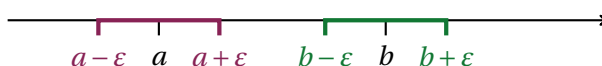
alternierend

²Beachten Sie, dass solche Definitionen auch für reelle Folgen gelten, weil alle reellen Folgen natürlich komplexe Folgen sind.

Sie sich die Definition noch einmal an, um zu sehen, dass dieser Fall dort nicht „verboten“ wurde.)

Das letzte Beispiel lässt sich offenbar verallgemeinern: Ist a irgendeine komplexe Zahl, dann konvergiert die konstante Folge $(a)_{n \in \mathbb{N}}$ gegen a .

Wenn eine Folge gegen eine Zahl a konvergiert, dann kann sie nicht gleichzeitig gegen eine *andere* Zahl b konvergieren. Das zeigt die folgende Skizze.



Man könnte nämlich dann ε so klein wählen, dass die Intervalle (bzw. Kreise) disjunkt sind, z. B. $\varepsilon = |b - a|/2$. Es können offensichtlich nicht einerseits fast alle Folgenwerte im roten und andererseits fast alle Folgenwerte im grünen Intervall liegen. Aus diesem Grund ergibt die folgende Definition Sinn.

Definition

Konvergiert eine Folge $(a_n)_{n \in \mathbb{N}}$ gegen eine Zahl a , so bezeichnet man diese Zahl als **den Grenzwert** oder **Limes** der Folge und schreibt

$$\lim_{n \rightarrow \infty} a_n = a.$$

Eine Folge, deren Grenzwert null ist, nennt man auch **Nullfolge**.

Das erste obige Beispiel hat also $\lim_{n \rightarrow \infty} 1/n = 0$ gezeigt.

Man kann sich nun einfache „Rechenregeln“ für Folgen herleiten, die wir an dieser Stelle einfach notieren.³

Lemma 30.1

Seien (a_n) und (b_n) konvergente komplexe Zahlenfolgen mit den Grenzwerten a und b . Dann konvergieren auch die Folgen $(a_n + b_n)$, $(a_n - b_n)$ und $(a_n \cdot b_n)$, und zwar gegen $a + b$, $a - b$ und ab . Gilt $b \neq 0$, dann konvergiert auch (a_n/b_n) , und zwar gegen a/b .

Beispielsweise konvergiert die Folge⁴ $((n+1)/n)$, die man auch als $(1 + 1/n)$ darstellen kann, gegen $1 + 0$, also gegen 1. Und $(1/n^2)$ konvergiert gegen $0 \cdot 0 = 0$, weil man das als $(1/n \cdot 1/n)$ interpretieren kann.

Aufgabe 30.6. Geben Sie mithilfe der obigen Rechenregeln Grenzwerte für $(1/n^3)$, $(5/n)$ und $(2n - 1)/n$ an.

Aufgabe 30.7. Konvergiert die Folge $((-1)^n)$? Und wenn ja, gegen welchen Grenzwert?

★Aufgabe 30.8. Übrigens kann es sein, dass in der Folge (a_n/b_n) in Lemma 30.1 einige Terme nicht definiert sind, weil evtl. durch null dividiert wird. Begründen Sie, dass das jedoch nur endlich viele sein können, wenn (b_n) nicht gegen null geht.

³Einen Beweis finden Sie bei Interesse im [alten Skript](#).

⁴Wie weiter oben angekündigt, wird ab jetzt der Zusatz „ $n \in \mathbb{N}$ “ meistens weggelassen, wenn das nicht zu Missverständnissen führen kann.

Für die Folge $((2n+5)/(3n+2))$ kann man einen „Trick“ anwenden: man multipliziert Zähler und Nenner mit $1/n$. (Das ist legitimes **Erweitern von Brüchen** und nur problematisch für den uninteressanten Fall $n = 0$.) Mit mehrfacher Anwendung der Rechenregeln (erst getrennt oben und unten die Regel für Summen, dann die Regel für die Division) erhält man

$$\frac{2n+5}{3n+2} = \frac{2n+5}{3n+2} \cdot \frac{1/n}{1/n} = \frac{2+5/n}{3+2/n} \rightarrow \frac{2+5 \cdot 0}{3+2 \cdot 0} = \frac{2}{3}.$$

Dabei ist zu beachten, dass wir die Regel für den Quotienten nur anwenden durften, weil der Nenner nicht gegen null konvergierte.

Aufgabe 30.9. Berechnen Sie die Grenzwerte der Folgen

$$\left(\frac{3n^2 + 5n - 10}{3 - 4n - 6n^2} \right) \quad \text{und} \quad \left(\frac{3n^2 + 5n - 10}{3 - 4n - 6n^3} \right).$$

(Tipp: Verwenden Sie denselben „Trick“ wie eben, wobei links mit n^{-2} und rechts mit n^{-3} erweitert wird.)

Aufgabe 30.10. Welchen Wert muss $a \in \mathbb{R}$ haben, damit diese Folge den Grenzwert 42 hat?

$$\left(\frac{3 + 4n + 5n^2 + 6n^3}{10 - 7n + an^3} \right)_{n \in \mathbb{N}}$$

Aufgabe 30.11. Geben Sie die kleinste natürliche Zahl k an, die dafür sorgt, dass diese Folge eine Nullfolge ist.

$$\left(\frac{(4n^3 + 3) \cdot (3n^4 + 4)}{12n^{2k} + 7} \right)_{n \in \mathbb{N}}$$

Im folgenden Satz werden die Grenzwerte von ein paar Folgen, die man kennen sollte, zusammengefasst. Der Vollständigkeit halber kommt auch ein Grenzwert vor, den wir bereits kennen.

Wichtige Folgen

- (i) Die **harmonische Folge** $(1/n)$ konvergiert gegen 0.
- (ii) Die **geometrische Folge** (q^n) konvergiert gegen 0, wenn q eine komplexe Zahl mit $|q| < 1$ ist. Gilt $|q| \geq 1$, so konvergiert diese Folge *nicht*.
- (iii) Für jede positive reelle Zahl a konvergiert $(\sqrt[n]{a})$ gegen 1.
- (iv) $(\sqrt[n]{n})$ konvergiert ebenfalls gegen 1.
- (v) Für $k \in \mathbb{N}$ und $b \in (1, \infty)$ ist (n^k/b^n) eine Nullfolge.

Satz 30.2

Die erste Folge hat etwas mit der **Harmonik** in der Musik zu tun, weil die Wellenlängen der **Obertöne** eines Grundtons $1/2$, $1/3$ und so weiter sind. Die zweite Folge heißt so, weil jedes Folgenglied das **geometrische Mittel** seiner beiden Nachbarn ist. Warum man wiederum diesen Mittelwert **geometrisch** nennt, wird in dem verlinkten Wikipedia-Artikel anhand von Beispielen ganz gut erklärt. Die letzte Folge im

obigen Satz spielt eine wichtige Rolle in der Informatik und wird insbesondere im **nächsten Abschnitt** wieder vorkommen.

Aufgabe 30.12. In einer konvergenten geometrischen Folge ist jedes Folgenglied kleiner als das vorherige. (In der Folge $(3/5)^n$ nimmt man z. B. in jedem Schritt 40 % vom vorherigen Wert weg.) Ist es deshalb nicht offensichtlich, dass so eine Folge gegen null gehen muss? Nein! Geben Sie ein Beispiel für eine Folge von positiven Zahlen an, bei der jedes Folgenglied kleiner als das vorherige ist, obwohl die Folge *nicht* gegen null geht.

Aufgabe 30.13. Berechnen Sie den Grenzwert der Folge $((4^{n+1} + 2^n)/(4^n + 2^{2n}))$. Wenden Sie dabei einen ähnlichen „Trick“ wie in Aufgabe Aufgabe 30.9 an.

Aufgabe 30.14. Berechnen Sie den Grenzwert der Folge $(\sqrt[n]{2^{n+1}})$.

Aufgabe 30.15. Wählen Sie in der Folge (n^k/b^n) die Werte $k = 10$ und $b = 1.001$ und berechnen Sie mithilfe eines Computers die ersten Folgenglieder. (Rechnen Sie durchgehend mit **Fließkommazahlen** oder verwenden Sie eine Programmiersprache, in der es keinen **Überlauf** bei ganzzahligen Berechnungen gibt.) Was fällt Ihnen auf?

Aufgabe 30.16. In **Wolfram Alpha** kann man sich mit Eingaben wie

limit of $(3n^3+2n-100)/(n^2-5n^3)$ as n to infinity

Grenzwerte berechnen lassen. Das kann man verwenden, um sich selbst Aufgaben zu stellen oder eigene Lösungen zu überprüfen. Man sollte nur jeweils zuerst nachschauen, ob das System die Eingabe so verstanden hat, wie sie intendiert war.

Definition

Eine Folge (a_n) komplexer Zahlen nennen wir **beschränkt**, wenn die Menge $\{|a_n| : n \in \mathbb{N}\}$ **beschränkt** ist.

Eine Folge ist also beschränkt, wenn es einen Kreis mit dem Mittelpunkt 0 gibt, in dem alle Folgenglieder liegen. (Ist die Folge reell, dann wird aus dem Kreis ein Intervall der Form $(-a, a)$.) Die Folge (42) ist z. B. trivialerweise beschränkt, weil die relevante Menge der Beträge der Folgenglieder die endliche Menge $\{42\}$ ist. Interessanter ist jedoch die folgende Aussage.

Lemma 30.3

Jede konvergente Folge ist beschränkt.

Das liegt daran, dass wegen der Konvergenz fast alle Folgenglieder in einem Kreis mit dem Radius 1 um den Grenzwert a liegen.⁵ Damit liegen fast alle Folgenglieder in einem Kreis mit dem Mittelpunkt 0 mit dem Radius $|a| + 1$. Sollte es Folgenglieder geben, die nicht in diesem Kreis liegen, so sind es nur endlich viele und man kann den Kreis schrittweise vergrößern, bis alle „eingefangen“ sind. (Machen Sie sich ggf. eine Skizze!)

⁵Dass wir gerade $\varepsilon = 1$ wählen, spielt keine Rolle. Man hätte jede positive Zahl nehmen können.

Aufgabe 30.17. Begründen Sie mit einem Gegenbeispiel, dass die Umkehrung von Lemma 30.3 nicht gilt, dass es also eine beschränkte Folge gibt, die nicht konvergiert.

Eine Folge, die nicht konvergiert, hatten wir *divergent* genannt. Dazu gehört das Verb *divergieren*. Für reelle Folgen gibt es allerdings noch eine spezielle Form von Divergenz, die wir nun definieren.

Sei (a_n) eine reelle Folge. Wenn es für alle $C > 0$ ein $n_0 \in \mathbb{N}$ mit $a_n \geq C$ für alle $n \geq n_0$ gibt, dann sagt man, dass die Folge **bestimmt gegen ∞ divergiert**. Man sagt ferner, dass eine reelle Folge (b_n) **bestimmt gegen $-\infty$ divergiert**, wenn $(-b_n)$ bestimmt gegen ∞ divergiert.

Definition

Wenn (a_n) bestimmt gegen ∞ divergiert, dann verwendet man Schreibweisen wie $\lim_{n \rightarrow \infty} a_n = \infty$ oder $a_n \rightarrow \infty$ und sagt auch manchmal, dass die Folge *gegen ∞ geht*. Dabei darf jedoch nicht vergessen werden, dass erstens ∞ *keine reelle Zahl ist* und dass zweitens so eine Folge *nicht* konvergiert und ∞ *kein* Grenzwert ist.

Das Standardbeispiel für eine Folge, die gegen ∞ divergiert, ist die Folge (n) . Und offensichtlich gilt das auch für Folgen wie (n^2) oder (2^n) . Die Gültigkeit der folgenden Aussage sollte hoffentlich klar sein.

Sei (a_n) eine reelle Folge, die bestimmt divergiert. Dann ist diese Folge unbeschränkt und $(1/a_n)$ ist eine Nullfolge.

Lemma 30.4

Aufgabe 30.18. Geben Sie ein Beispiel für eine reelle Nullfolge (a_n) an, für die $(1/a_n)$ nicht bestimmt divergiert.⁶

Bemerkungen

Beachten Sie, dass n in einem Ausdruck wie $(2^n)_{n \in \mathbb{N}}$ eine sogenannte *gebundene Variable* ist, die man durch eine andere ersetzen kann. $(2^k)_{k \in \mathbb{N}}$ ist *dieselbe* Folge! Siehe die Anmerkung zu gebundenen Variablen auf Seite 31.

Es sei noch einmal darauf hingewiesen, dass ∞ und $-\infty$ *keine Zahlen* sind und daher trotz der suggestiven Schreibweise $\lim_{n \rightarrow \infty} a_n = \infty$ keine Grenzwerte sein können. Eine bestimmt divergierende Folge konvergiert *nicht*!

Außerdem ergibt eine Aussage wie $\lim_{n \rightarrow \infty} a_n = \infty$ für beliebige Folgen *komplexer* Zahlen keinen Sinn, weil man in $\mathbb{C}$ Zahlen nicht mit $<$ vergleichen kann.

Wenn im Skript die Formulierung *fast alle* verwendet wird, dann ist damit das oben definierte Konzept gemeint und nicht etwa eine vage Vorstellung, über die man geteilter Meinung sein könnte. Die mathematische Fachsprache „kapert“ sich manchmal Begriffe aus der Alltagssprache und gibt ihnen eine neue Bedeutung. Siehe dazu auch die Bemerkungen zu Definitionen auf Seite 6.

⁶ $(1/a_n)$ soll definiert sein, d. h. fast alle Folgenglieder von (a_n) müssen von null verschieden sein.

31. Die Landau-Notation

In diesem kurzen Abschnitt führen wir eine Schreibweise ein, die in der Informatik häufig benutzt wird, um die **Komplexität** von Algorithmen zu beschreiben. Sie ist benannt nach dem deutschen Mathematiker **Edmund Landau** (obwohl sie eigentlich von Landaus Landsmann **Paul Bachmann** eingeführt wurde).

Aufgabe 31.1. Wie viele Rechenoperationen (Multiplikationen und Additionen) muss man durchführen, wenn man zwei $n \times n$ -Matrizen multipliziert?

Wenn bei dem Ausdruck aus der Lösung von Aufgabe 31.1 n wächst, dann verhält er sich mehr und mehr wie $2n^3$, weil für große n der Term n^2 kaum noch eine Rolle spielt. Außerdem unterscheiden sich $2n^3$ und n^3 nur durch den konstanten Faktor 2. In der Informatik (und nicht nur da) gibt es viele Situationen, in denen man sagen würde, dass die Anzahl der benötigten Operationen bei wachsendem n „ungefähr wie n^3 “ wächst; man möchte gewissermaßen „unwichtige Informationen“ weglassen. Das wird durch die folgende Definition präzisiert.

Definition

Sind f und g Folgen reeller Zahlen, so sagt man, dass g **von der Ordnung f** oder auch „**groß-Oh**“ **von f** ist,⁷ wenn es ein $C > 0$ mit $|g(n)| \leq C|f(n)|$ für alle $n \in \mathbb{N}$ gibt. Für die Menge aller reellen Folgen g mit dieser Eigenschaft schreibt man $\mathcal{O}(f)$.

$g \in \mathcal{O}(f)$ ist also gleichbedeutend damit, dass g von der Ordnung f ist.

In der Informatik sind die Folgen, für die diese Notation verwendet wird, in der Regel solche, die bestimmt gegen ∞ konvergieren, z. B. (2^n) . Gilt $g \in \mathcal{O}(f)$, so heißt das, dass g *nicht wesentlich stärker anwächst* als f , denn es gibt ja eine feste Zahl C , für die immer $g(n) \leq C f(n)$ gilt. Die Folgenglieder von g sind also höchstens C -mal so groß wie die von f zum selben Index. Entsprechend bedeutet $g \notin \mathcal{O}(f)$, dass g deutlich schneller wächst als f .

Konvention

Wenn wir in diesem Abschnitt Schreibweisen wie n^2 verwenden, dann sind damit häufig die entsprechenden Folgen gemeint. Statt $3n^2 \in \mathcal{O}(n^2)$ müsste man eigentlich $(3n^2)_{n \in \mathbb{N}} \in \mathcal{O}((n^2)_{n \in \mathbb{N}})$ schreiben, aber das wäre natürlich sehr umständlich.

Das Beispiel $3n^2 \in \mathcal{O}(n^2)$ ist korrekt, denn wir können hier $C = 3$ wählen. Andererseits gilt auch $n^2 \in \mathcal{O}(3n^2)$, weil wir in diesem Fall $C = 1/3$ verwenden können. Der Sinn der Landau-Notation ist allerdings wie gesagt das Vereinfachen der Terme bzw. das Weglassen „unwichtiger“ Anteile. Daher ist die erste Variante, bei der der Faktor 3 „entfernt“ wurde, die typische Verwendung dieser Schreibweise. Für das Beispiel hinter Aufgabe 31.1 gilt wie bereits angedeutet $2n^3 - n^2 \in \mathcal{O}(n^3)$, denn offenbar gilt immer $|2n^3 - n^2| \leq 1/2 \cdot |n^3|$.

⁷Im Englischen spricht man deshalb auch von der *big O notation*.

Zum „Rechnen“ mit der Landau-Notation folgt nun die zentrale Aussagen.

Seien f_1, f_2, g_1, g_2 reelle Folgen mit $g_i \in \mathcal{O}(f_i)$ für $i = 1, 2$. Dann gilt:

- (i) $\lambda g_1 \in \mathcal{O}(f_1)$ für $\lambda \in \mathbb{R}$
- (ii) $g_1 + g_2 \in \mathcal{O}(f_1)$, wenn $|f_2(n)| \leq |f_1(n)|$ für fast alle $n \in \mathbb{N}$ gilt.
- (iii) $g_1 g_2 \in \mathcal{O}(f_1 f_2)$
- (iv) $\mathcal{O}(g_1) \subseteq \mathcal{O}(f_1)$

Satz 31.1

Betrachtet man Folgen wie die aus Aufgabe 31.1 als Angaben über den Zeitaufwand für einen Algorithmus in Abhängigkeit von der Eingabegröße n , so kann man (ii) so interpretieren, dass bei aufeinanderfolgenden Abschnitten des Verfahrens nur der langsamste relevant ist. (iii) besagt, dass in Schleifen die (Ordnung der) Anzahl der Schleifendurchläufe mit dem Aufwand für einen Durchlauf des Schleifenkörpers multipliziert werden muss. (i) ist schließlich ein Zeichen dafür, dass dieses Maß für Aufwände relativ grob ist, weil konstante Faktoren einfach ignoriert werden. Man könnte es etwas salopp auch so formulieren, dass man einen Unterschied wie n^3 zu $2n^3$ einfach durch den Kauf eines doppelt so schnellen Computers wettmachen könnte.

Ein wichtiges Hilfsmittel, um zu entscheiden, ob eine Aussage wie $a_n \in \mathcal{O}(b_n)$ gilt, ist das folgende Lemma.

Sind f und g Folgen von reellen Zahlen und ist $(g(n)/f(n))$ definiert⁸ und konvergent, dann gilt $g \in \mathcal{O}(f)$. Ist diese Folge sogar eine Nullfolge, dann gilt $f \notin \mathcal{O}(g)$.

Lemma 31.2

Man erhält damit insbesondere:

Seien $k, m \in \mathbb{N}$ und $a, b \in (1, \infty)$. Dann gilt:

- (i) Aus $k < m$ folgt $\mathcal{O}(n^k) \subset \mathcal{O}(n^m)$.
- (ii) Aus $a < b$ folgt $\mathcal{O}(a^n) \subset \mathcal{O}(b^n)$.
- (iii) $\mathcal{O}(n^k) \subset \mathcal{O}(a^n)$

Korollar 31.3

Mit den obigen Aussagen ergibt sich z. B. $5 \in \mathcal{O}(n^2)$ wegen $5 = 5 \cdot n^0$ und $-2n \in \mathcal{O}(n^2)$ und $8n^2 \in \mathcal{O}(n^2)$ sowie schließlich $8n^2 - 2n + 5 \in \mathcal{O}(n^2)$ für die Summe. Wie oben bereits erwähnt ist es der „Job“ der Landau-Notation, die für das asymptotische Wachstumsverhalten „unwichtigen“ Anteile zu entfernen. Mathematisch wäre natürlich auch $8n^2 - 2n + 5 \in \mathcal{O}(n^3)$ korrekt, aber in den Anwendungen ist man immer an einer möglichst „scharfen“ Schranke interessiert.

Korollar 31.3 beschreibt die wesentlichen „Stufen“, die in der Informatik eine Rolle spielen. Außer denen und deren Beziehungen untereinander, braucht man

⁸Damit ist gemeint, dass fast alle Folgenglieder von f von null verschieden sind, damit nicht zu oft durch null dividiert wird.

manchmal noch eine weitere, die hier der Vollständigkeit halber schon einmal erwähnt sein soll, obwohl wir erst später über [Logarithmen](#) sprechen werden.: Es gilt $\mathcal{O}(\log_a n) \subset \mathcal{O}(n)$ für alle $a > 0$. Da außerdem $\mathcal{O}(\log_a n) = \mathcal{O}(\log_b n)$ für alle $a, b > 0$ gilt, schreibt man in der Regel einfach $\mathcal{O}(\log n)$.

Aufgabe 31.2. Welche der folgenden Aussagen sind wahr?

- (i) $2^{n+42} \in \mathcal{O}(3^n)$
- (ii) $3^n \in \mathcal{O}(2^{n+42})$
- (iii) $n^2 2^n \in \mathcal{O}(2^n)$
- (iv) $n^2 2^n \in \mathcal{O}(3^n)$
- (v) $(2^n)^2 \in \mathcal{O}(2^n)$
- (vi) $n \log n \in \mathcal{O}(n^2)$

Bemerkungen

In Fachtexten findet man häufig Ausdrücke wie $n^2 + 5n + 2 = \mathcal{O}(n^2)$, d. h., statt $\in$ wird $=$ verwendet. Das ist leider didaktisch sehr unglücklich, weil es sich *nicht* um eine Gleichung handelt. Wenn Sie so etwas lesen, dann sollten Sie es im Geiste immer in die korrekte Schreibweise aus diesem Abschnitt „übersetzen“.

★ Die Landau-Notation im Computer

Die in diesem Skript bereits erwähnte Bibliothek SYMPY „versteht“ auch die Landau-Notation. Details entnehmen Sie bitte der [Dokumentation](#).

```
from sympy import *
n = symbols("n")

O(8*n**2-2*n+5, (n, oo))
2**(n+42) in O(3**n, (n, oo))
```

32. Reihen

Reihen sind so etwas wie „unendliche Summen“, allerdings kann man den Begriff der Summe nicht ohne Weiteres auf unendlich viele Summanden erweitern. Zunächst führen wir eine Schreibweise für endliche Summen ein, die manchen aus der Schule bekannt sein wird.

Definition

Seien m, n natürliche Zahlen und $(a_k)_{k \in \mathbb{N}}$ eine komplexe Folge. Für $m \leq n$ definieren wir

$$\sum_{k=m}^n a_k = a_m + a_{m+1} + a_{m+2} + \cdots + a_n \quad (32.1)$$

und für $m > n$ soll $\sum_{k=m}^n a_k$ null sein.

Offensichtlich müssen nicht alle Folgenglieder der Folge definiert sein. Es reicht, wenn a_m bis a_n sinnvolle Ausdrücke sind. Beachten Sie, dass k als sogenannte *Laufvariable* gebunden und daher in gewissem Sinne bedeutungslos ist. (Siehe die entsprechende [Bemerkung zu Folgen](#).) Der Wert $\sum_{k=m}^n a_k$ ändert sich also nicht, wenn man stattdessen $\sum_{i=m}^n a_i$ schreibt.

Falls Sie Schwierigkeiten mit dieser Schreibweise haben, schauen Sie sich den [Abschnitt über Summen im Computer](#) an.

Zur Schreibweise ist außerdem zu sagen, dass nach wie vor „Punkt- vor Strichrechnung“ gilt, so dass man bei einem Ausdruck wie $\sum_{i=2}^8 (5i)$ die Klammern in der Regel weglassen wird, weil die Multiplikation $5i$ eine höhere Priorität als die Summe hat. Bei $\sum_{i=2}^8 (i + 5)$ gilt das jedoch nicht, weil man diese Summe ohne Klammern für $5 + \sum_{i=2}^8 i$ halten könnte.

Aufgabe 32.1. Geben Sie den Wert $\sum_{k=3}^5 k^2$ an.

Aufgabe 32.2. Geben Sie die Summen

$$\begin{aligned} s_1 &= 4 + 8 + 12 + 16 + \cdots + 404, \\ s_2 &= 10 + 13 + 16 + 19 + \cdots + 3001, \\ s_3 &= 1 + 1/4 + 1/9 + 1/16 + 1/25 + \cdots + 1/121, \\ s_4 &= 1 - 2 + 3 - 4 + 5 - 6 + \cdots + 41 - 42 \quad \text{und} \\ s_5 &= na_6 + na_9 + na_{12} + \cdots + na_{99} \end{aligned}$$

mithilfe der eben eingeführten Schreibweise an.

Da es sich um ganz normale Addition handelt, kann man bekannte Rechenregeln anwenden, zum Beispiel die im nächsten Lemma zusammengefassten.

Für endliche Summen gelten die folgenden Rechenregeln.

Lemma 32.1

$$\begin{aligned} \sum_{k=m}^n (a_k \pm b_k) &= \left(\sum_{k=m}^n a_k \right) \pm \sum_{k=m}^n b_k \\ \sum_{k=m}^n \lambda a_k &= \lambda \sum_{k=m}^n a_k && \text{für } \lambda \in \mathbb{C} \\ \sum_{k=m}^n a_k &= \left(\sum_{k=0}^n a_k \right) - \sum_{k=0}^{m-1} a_k && \text{für } n \geq m \geq 0 \\ \sum_{k=m}^n a_k &= \sum_{k=0}^{n-m} a_{k+m} \end{aligned}$$

Aufgabe 32.3. Machen Sie sich anhand von Beispielen klar, was die Regeln in Lemma 32.1 besagen und warum sie korrekt sind.

Aufgabe 32.4. Begründen Sie mit einem einfachen Gegenbeispiel, dass die erste Regel aus Lemma 32.1 für die Multiplikation nicht stimmt, dass

$$\sum_{k=m}^n a_k b_k = \left(\sum_{k=m}^n a_k \right) \cdot \sum_{k=m}^n b_k$$

also im Allgemeinen *nicht* korrekt ist.

Für bestimmte regelmäßige Summen kann man sich das Addieren sparen, weil eine einfache Formel für den Wert bekannt ist. Die beiden wichtigsten sind die folgenden.

Satz 32.2

Gaußsche (bzw. arithmetische) und geometrische Summenformel

Für alle $n \in \mathbb{N}$ und alle $q \in \mathbb{C} \setminus \{1\}$ gilt:

$$\sum_{k=1}^n k = \frac{n(n+1)}{2}$$

$$\sum_{k=0}^n q^k = \frac{q^{n+1} - 1}{q - 1}$$

Als Begründung für die erste Formel schreibt man die Summe zweimal auf – einmal vorwärts und darunter noch einmal rückwärts – und addiert dann alle Terme:

$$\begin{array}{cccccccc} 1 & 2 & 3 & \cdots & (n-2) & (n-1) & n \\ n & (n-1) & (n-2) & \cdots & 3 & 2 & 1 \\ \hline (n+1) & (n+1) & (n+1) & \cdots & (n+1) & (n+1) & (n+1) \end{array}$$

Man beobachtet, dass paarweise je ein Summand oben und unten zusammen immer $n+1$ ergeben, so dass bei insgesamt n Summanden $n \cdot (n+1)$ als Gesamtsumme herauskommt. Weil man die Summe zweimal hingeschrieben hat, muss man am Ende noch durch 2 dividieren.

Für die zweite Formel setzen wir $S = \sum_{k=0}^n q^k$ und „verschieben“ alle Summanden durch Multiplikation mit dem Faktor q :

$$\begin{aligned} (q-1)S &= qS - S = \left(\sum_{k=0}^n q^{k+1} \right) - \sum_{k=0}^n q^k \\ &= q^{n+1} + q^n + \cdots + q^2 + q^1 - (q^n + q^{n-1} + \cdots + q^1 + q^0) \\ &= q^{n+1} - q^0 = q^{n+1} - 1. \end{aligned}$$

Dabei ergibt sich die vorletzte Gleichheit, weil q^1, q^2 und so weiter bis q^n jeweils links addiert und dann rechts wieder subtrahiert werden. Es bleiben nur zwei Terme übrig. Dividiert man nun auf beiden Seiten durch $q-1$ (was OK ist, weil nach Voraussetzung $q \neq 1$ gilt), so ist man fertig.

Aufgabe 32.5. Verwenden Sie Lemma 32.1 und Satz 32.2 zusammen, um die Summen $\sum_{n=5}^{20} 3n$ sowie $\sum_{n=3}^8 1/2^n$ zu berechnen.

Aufgabe 32.6. Schreiben Sie $40 + 50 + 60 + \dots + 230$ mittels des Summenzeichens auf und berechnen Sie dann diese Summe.

Aufgabe 32.7. Geben Sie die Summe $7 + 10 + 13 + 16 + \dots + 1000$ an.

Aufgabe 32.8. Es gibt außer den beiden in Satz 32.2 aufgeführten noch weitere Summenformeln. Wichtig ist z. B. die **Faulhabersche Formel**. Schauen Sie sich den verlinkten Wikipedia-Artikel an und berechnen Sie mithilfe dieser Formel die Summen $\sum_{n=1}^8 n^2$ sowie $3^2 + 4^2 + 5^2 + \dots + 30^2$.

Jetzt können wir mit dem eigentlichen Thema beginnen: mit den „unendlichen Summen“. Zunächst sollte man sich klarmachen, dass es so etwas gar nicht gibt, wenn man es nicht definiert, weil Addition nur für zwei Summanden (und damit für endlich viele) definiert ist. Dass eine naive Herangehensweise nicht funktioniert, zeigt das folgende Beispiel:

$$\begin{aligned} 1 + (-1) + 1 + (-1) + 1 + (-1) + \dots &= (1 - 1) + (1 - 1) + (1 - 1) + \dots = 0 \\ 1 + (-1) + 1 + (-1) + 1 + (-1) + \dots &= 1 - (1 - 1) - (1 - 1) - (1 - 1) - \dots = 1 \\ 1 + (-1) + 1 + (-1) + 1 + (-1) + \dots &= (-1) + 1 + (-1) + 1 + (-1) + 1 + \dots \\ &= -1 + (1 - 1) + (1 - 1) + (1 - 1) + \dots = -1 \end{aligned}$$

Was kommt denn da nun heraus? 0, 1, -1 oder etwas ganz anderes? Es sieht jedenfalls so aus, als könne man nicht einfach „drauflos addieren“ und bekannte Regeln wie **Assoziativität und Kommutativität** ins Unendliche „übersetzen“.

Sei $(a_k)_{k \in \mathbb{N}}$ eine komplexe Folge und $m \in \mathbb{N}$. Für die Folge $(\sum_{k=m}^n a_k)_{n \in \mathbb{N}}$ schreiben wir $\sum_{k=m}^{\infty} a_k$ und sprechen von der **Reihe** mit den **Reihengliedern** a_k und den **Partialsummen** $\sum_{k=m}^n a_k$. Wenn die Reihe (also die Folge der Partialsummen) konvergiert, so nennt man den Grenzwert den **Wert** bzw. die **Summe** der Reihe und schreibt dafür $\sum_{k=m}^{\infty} a_k$.

Definition

$\sum_{k=m}^{\infty} a_k$ bezeichnet also einerseits die Reihe und andererseits den Wert der Reihe (wenn er existiert). Man kann somit sagen, dass $\sum_{k=m}^{\infty} a_k$ konvergiert, andererseits ergibt aber auch so etwas wie $\sum_{k=m}^{\infty} a_k = 42$ Sinn. Das ist zwar theoretisch missverständlich, aber historisch so gewachsen und führt erfahrungsgemäß kaum zu Problemen.

Eher problematisch ist anfangs für viele, dass in der Definition zwei Folgen vorkommen, die man auseinanderhalten muss. Geht es beispielsweise um $a_k = 1/2^k$, so haben wir erstens die Folge $1, 1/2, 1/4, 1/8, \dots$ und zweitens die Folge

$$\begin{aligned} &1, \\ &1 + 1/2 = 3/2, \\ &1 + 1/2 + 1/4 = 7/4, \\ &1 + 1/2 + 1/4 + 1/8 = 15/8 \quad \text{und so weiter.} \end{aligned}$$

Die Idee ist, dass $\sum_{k=0}^{\infty} 1/2^k$ so etwas wie die „Summe“

$$1 + 1/2 + 1/4 + 1/8 + 1/16 + \dots$$

darstellen soll, wobei die erste Folge die Folge der „Summanden“ (Reihenglieder) ist, während die zweite die der „Zwischenergebnisse“ (Partialsummen) ist. Die Begriffe stehen in Häkchen, weil es keine unendlichen Summen gibt, solange man nur Addition hat. Erst durch das Konzept der Reihe (für das man Folgen und Grenzwerte braucht) wird klar definiert, was damit gemeint sein soll.

Außerdem noch eine Klarstellung zum „Anfang“ einer Reihe:

Lemma 32.3

Sei (a_n) eine komplexe Folge und $m \in \mathbb{N}$. Die Reihe $\sum_{n=m}^{\infty} a_n$ konvergiert genau dann, wenn $\sum_{n=0}^{\infty} a_n$ konvergiert. Es gilt dann $\sum_{n=0}^{\infty} a_n = \left(\sum_{n=0}^{m-1} a_n\right) + \sum_{n=m}^{\infty} a_n$.

Das folgt direkt aus der Definition, denn die eine Reihe ist die Folge $(\sum_{k=m}^n a_k)$, die andere ist $(\sum_{k=0}^n a_k)$, d. h., die zweite Folge erhält man aus der ersten, indem man zu jedem Folgenglied $\sum_{n=0}^{m-1} a_n$ addiert. Für die Konvergenz ist dieses Addieren eines konstanten Wertes egal, für den Grenzwert natürlich nicht. Außerdem kann man sicher den „Start“ $n = 0$ auch durch $n = 1$ oder einen anderen Wert ersetzen, wenn man das an beiden Stellen macht.

Aufgabe 32.9. Was ist mit dem Beispiel $1 + (-1) + 1 + (-1) + \dots$ von oben? Wie würde man es als Reihe schreiben und was für ein Ergebnis bekommt man?

Ein wichtiges notwendiges Kriterium für die Konvergenz einer Reihe ist der folgende Satz:

Satz 32.4

Nullfolgenkriterium

Wenn $\sum_{k=0}^{\infty} a_k$ konvergiert, dann ist (a_n) eine Nullfolge.

Das Nullfolgenkriterium kann man sich anschaulich folgendermaßen klarmachen: Wenn die Reihe gegen die Zahl S konvergiert, dann liegen nach Definition der Konvergenz zu vorgegebenem $\varepsilon > 0$ fast alle Partialsummen im Intervall $(S - \varepsilon, S + \varepsilon)$. Die zugehörigen Summanden (Reihenglieder) müssen dann aber auf jeden Fall einen kleineren Betrag als 2ε haben, weil das Intervall sonst beim Übergang von einer Partialsumme s_n zur nächsten verlassen würde. Weil ε beliebig klein werden kann, müssen auch die Beträge der Summanden beliebig klein werden.⁹



harmonische Reihe

Man beachte, dass Satz 32.4 lediglich ein notwendiges, aber kein hinreichendes Kriterium ist. Dafür schauen wir uns als wichtiges Beispiel die *harmonische Reihe*

⁹Stellen Sie sich zur Übung vor, wie die folgende Skizze im Komplexen aussehen müsste. Dort gilt das Nullfolgenkriterium ja auch.

$\sum_{n=1}^{\infty} 1/n$ an. Nach einer Beweisidee von **Nikolaus von Oresme** aus dem 14. Jahrhundert betrachtet man dafür ausgewählte Partialsummen nach dem folgenden Muster:

$$\sum_{k=1}^1 \frac{1}{k} = \frac{1}{2} + \frac{1}{2} = 2 \cdot \frac{1}{2}$$

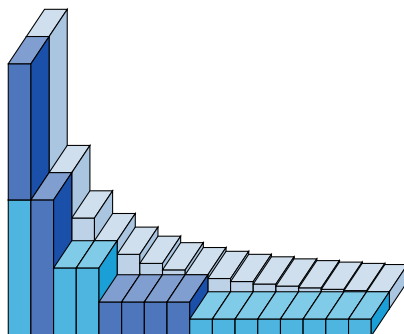
$$\sum_{k=1}^2 \frac{1}{k} = \sum_{k=1}^1 \frac{1}{k} + \frac{1}{2} = 3 \cdot \frac{1}{2}$$

$$\sum_{k=1}^4 \frac{1}{k} = \sum_{k=1}^2 \frac{1}{k} + \frac{1}{3} + \frac{1}{4} \geq 3 \cdot \frac{1}{2} + 2 \cdot \frac{1}{4} = 4 \cdot \frac{1}{2}$$

$$\sum_{k=1}^8 \frac{1}{k} = \sum_{k=1}^4 \frac{1}{k} + \frac{1}{5} + \frac{1}{6} + \frac{1}{7} + \frac{1}{8} \geq 4 \cdot \frac{1}{2} + 4 \cdot \frac{1}{8} = 5 \cdot \frac{1}{2}$$

$$\sum_{k=1}^{16} \frac{1}{k} = \sum_{k=1}^8 \frac{1}{k} + \frac{1}{9} + \frac{1}{10} + \dots + \frac{1}{16} \geq 5 \cdot \frac{1}{2} + 8 \cdot \frac{1}{16} = 6 \cdot \frac{1}{2}$$

Allgemein ist die 2^m -te Partialsumme offenbar immer mindestens so groß wie $(m+2)/2$. Daher ist die Folge der Partialsummen unbeschränkt und die Reihe konvergiert nicht, *obwohl* die Folge der Reihenglieder das Nullfolgenkriterium erfüllt!



Nun aber mal zwei Reihen, die tatsächlich konvergieren:

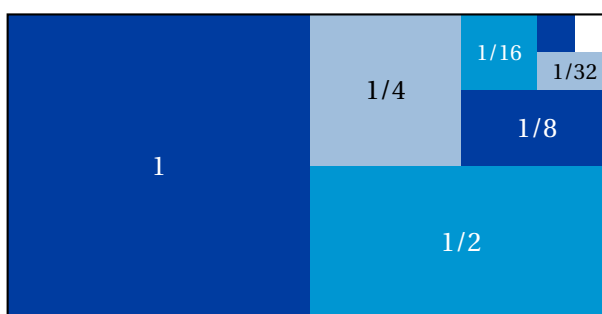
Geometrische Reihe

Sei $q \in \mathbb{C}$. Für $|q| < 1$ konvergiert $\sum_{n=0}^{\infty} q^n$ gegen $(1-q)^{-1}$. Für $|q| \geq 1$ divergiert diese Reihe.

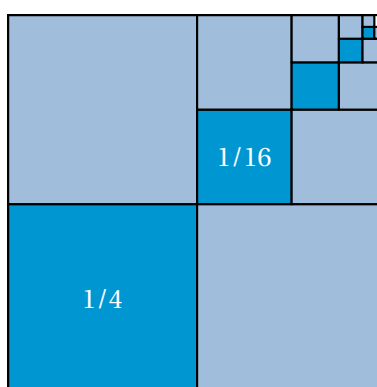
Satz 32.5

Nach Satz 32.2 ist die Folge der Partialsummen $((q^{n+1} - 1)/(q - 1))$. Da q^{n+1} nach Satz 30.2 für $|q| < 1$ eine Nullfolge ist, folgt sofort der erste Teil der Behauptung. Der zweite Teil ergibt sich aus dem **Nullfolgenkriterium**.

Für $q = 1/2$ ergibt sich 2 als Reihenwert. Das kann man an der folgenden Skizze ganz gut geometrisch sehen.



Und für $q = 1/4$ erhält man $\sum_{k=1}^{\infty} 4^{-k} = 1/3$. (Rechnen Sie es nach!) Das kann man aus der folgenden Grafik schließen, in der jeweils eines von drei gleich großen Quadraten (also ein Drittel) durch kräftiges Blau hervorgehoben ist.



Aufgabe 32.10. Wenden Sie Satz 32.5 und Lemma 32.3 an, um $\sum_{n=2}^{\infty} 3^{-n}$ zu berechnen.

Lemma 32.6

Die Reihe $\sum_{k=1}^{\infty} (k^2 + k)^{-1}$ konvergiert gegen 1.

Das kann man in diesem Fall an der Umformung

$$\frac{1}{k^2 + k} = \frac{(k+1) - k}{k(k+1)} = \frac{1}{k} - \frac{1}{k+1}$$

für ein einzelnes Reihenglied erkennen. Damit erhält man mit Lemma 32.1 für die Partialsummen

$$\begin{aligned} \sum_{k=1}^n \frac{1}{k^2 + k} &= \sum_{k=1}^n \left(\frac{1}{k} - \frac{1}{k+1} \right) = \left(\sum_{k=1}^n \frac{1}{k} \right) - \sum_{k=1}^n \frac{1}{k+1} \\ &= \left(1 + \frac{1}{2} + \frac{1}{3} + \cdots + \frac{1}{n-1} + \frac{1}{n} \right) - \left(\frac{1}{2} + \frac{1}{3} + \cdots + \frac{1}{n} + \frac{1}{n+1} \right) \\ &= 1 - \frac{1}{n+1} \end{aligned}$$

mit dem Grenzwert 1 für $n \rightarrow \infty$.

Die beiden Reihen aus Satz 32.5 und Lemma 32.6 gehören zu den wenigen, für die man vergleichsweise einfach Grenzwerte berechnen kann.

Natürlich gibt es auch für Reihen ein paar Rechenregeln. Die einfachsten folgen direkt aus denen für endliche Summen, angewendet auf die Partialsummen.

Seien (a_n) und (b_n) komplexe Folgen, λ eine komplexe Zahl und m eine natürliche Zahl. Dann gilt:

- (i) Konvergiert $\sum_{n=m}^{\infty} a_n$ gegen a , dann konvergiert auch $\sum_{n=m}^{\infty} \lambda a_n$, und zwar gegen λa .
- (ii) Konvergieren $\sum_{n=m}^{\infty} a_n$ und $\sum_{n=m}^{\infty} b_n$ gegen a bzw. b , dann konvergiert auch $\sum_{n=m}^{\infty} (a_n \pm b_n)$, und zwar gegen $a \pm b$.

Lemma 32.7

Aufgabe 32.11. Berechnen Sie

$$\sum_{n=3}^{\infty} \frac{3}{n+n^2}, \quad \sum_{n=2}^{\infty} \frac{5}{2^n} \quad \text{und} \quad \sum_{n=0}^{\infty} \frac{2 \cdot 4^n + 3^{n+1}}{6^n}.$$

Nun können wir auch die Lücke schließen, die bei der [Definition der reellen Zahlen](#) offen geblieben war: Wie kann man das Konzept von „unendlich vielen Nachkommastellen“ präzisieren? Durch Reihen! Die Ziffer k an der n -ten Stelle hinter dem Komma steht wie im endlichen Fall für $k \cdot 10^{-n}$, aber man benötigt eine Reihe, um diese Brüche alle zu addieren. Es folgen Beispiele, die zeigen, wie das gemeint ist.

$$\begin{aligned} 0.\bar{3} &= \frac{3}{10} + \frac{3}{100} + \cdots = \sum_{n=1}^{\infty} 3 \cdot 10^{-n} = 3 \sum_{n=1}^{\infty} 10^{-n} \\ &= 3 \cdot \left(\frac{1}{1-1/10} - 1 \right) = 3 \cdot \frac{1}{9} = \frac{1}{3} \\ 0.\bar{9} &= \frac{9}{10} + \frac{9}{100} + \cdots = \sum_{n=1}^{\infty} 9 \cdot 10^{-n} = 9 \sum_{n=1}^{\infty} 10^{-n} \\ &= 9 \cdot \left(\frac{1}{1-1/10} - 1 \right) = 9 \cdot \frac{1}{9} = 1 \\ 0.\overline{27} &= \frac{27}{100} + \frac{27}{10000} + \cdots = \sum_{n=1}^{\infty} 27 \cdot 100^{-n} = 27 \sum_{n=1}^{\infty} 100^{-n} \\ &= 27 \cdot \left(\frac{1}{1-1/100} - 1 \right) = 27 \cdot \frac{1}{99} = \frac{3}{11} \\ 0.0\overline{27} &= \frac{27}{1000} + \frac{27}{100000} + \cdots = \sum_{n=1}^{\infty} 27 \cdot 10^{-(2n+1)} = \frac{1}{10} \cdot \sum_{n=1}^{\infty} 27 \cdot 100^{-n} = \frac{3}{110} \\ 0.4\overline{27} &= 0.4 + 0.0\overline{27} = \frac{2}{5} + \frac{3}{110} = \frac{47}{110} \end{aligned}$$

Es geht hier in allen Fällen um die [geometrische Reihe](#). Die Beispiele rechtfertigen im Nachhinein auch die Umrechnungstechniken, die ab Seite [94](#) entwickelt wurden.

★ Summen und Reihen im Computer

Für das Verständnis der Summenschreibweise ([32.1](#)) ist es evtl. sinnvoll, sich diese als Programm vorzustellen. In PYTHON würde es so aussehen:

```

S = 0
k = m
while k <= n:
    S += a(k)
    k += 1

```

Dabei wird vorausgesetzt, dass in den Variablen m und n natürliche Zahlen stehen und dass die PYTHON-Funktion a die Folge $(a_k)_{k \in \mathbb{N}}$ implementiert. Am Ende steht in S die Summe – und zwar auch dann, wenn m größer als n ist.

Für die Praxis bietet PYTHON mit `sum` allerdings bereits eine Funktion für das Berechnen endlicher Summen.

```
sum(k**2 for k in range(10))
```

Die bereits mehrfach erwähnte Bibliothek SYMPY kann auch mit Folgen und Reihen umgehen und deren Grenzwerte berechnen. Mehr dazu in *Konkrete Mathematik (nicht nur) für Informatiker* ab Kapitel 21. Und in Kapitel 18 wird gezeigt, wie man mit *Generatoren* unendliche Folgen simulieren kann.

33. Exponentialfunktion und Logarithmus

Die Exponentialfunktion modelliert eine in der Natur häufig vorkommende Form des Wachstums. Aus verschiedenen Gründen ist sie die wohl wichtigste Funktion in der Analysis.

Wir benötigen zunächst einen aus der Schule bekannten Hilfsbegriff:

Definition

Die **Fakultät** $n!$ einer natürlichen Zahl n ist das Produkt aller positiven natürlichen Zahlen, die nicht größer als n sind. Insbesondere soll $0!$ den Wert 1 haben.

Es gilt also z. B. $3! = 1 \cdot 2 \cdot 3$, $5! = 1 \cdot 2 \cdot 3 \cdot 4 \cdot 5$ und allgemein $(n+1)! = n \cdot n!$. Für $0!$ ergibt sich ein „leeres Produkt“, dessen Wert wie bei der *Definition der Summe* sinnvollerweise als das *neutrale Element* der Multiplikation festgesetzt wird.

$n!$ wächst extrem schnell, wenn n größer wird. Beispielsweise ist $10!$ größer als 3 Millionen und $20!$ größer als 2 *Millionen*. Ist $k \in \mathbb{N}$ eine fest gewählte Zahl, so „überholt“ $n!$ offensichtlich irgendwann den Wert k^n , wie man am folgenden Beispiel für $k = 4$ sieht.

n	0	1	2	3	4	5	6	7	8	9
4^n	1	4	16	64	256	1024	4096	16384	65536	262144
$n!$	1	1	2	6	24	120	720	5040	40320	362880

Das liegt daran, dass in jedem weiteren Schritt oben „nur“ mit 4 multipliziert wird, während unten der Faktor immer größer wird. Das halten wir im folgenden Lemma fest. (Man vergleiche mit der letzten Folge aus Satz 30.2.)

Für jedes $b \in (1, \infty)$ ist $(b^n/n!)$ eine Nullfolge.

Lemma 33.1

Aufgabe 33.1. Geben Sie den Wert von $n!/(n+1)!$ in Abhängigkeit von n an.

Aufgabe 33.2. Gilt $3^n \in \mathcal{O}(n!)$? Gilt $(n+1)! \in \mathcal{O}(n!)$?

Die durch

$$\exp(z) = \sum_{n=0}^{\infty} \frac{z^n}{n!} \quad (33.1)$$

für $z \in \mathbb{C}$ definierte Funktion wird (**natürliche**) **Exponentialfunktion** genannt. Für den Funktionswert $\exp(1)$ schreibt man e und nennt ihn die **eulersche Zahl**.

Definition

e ist übrigens eine **irrationale Zahl**, wie ihr Namenspatre **Leonhard Euler** 1737 bewies. Einen kurzen Beweis dafür finden Sie in dem Video *Die eulersche Zahl e ist irrational*.

Wichtig für die obige Definition ist natürlich, dass der Ausdruck in (33.1) für jede komplexe Zahl konvergiert, weil $\exp(z)$ sonst keine Zahl wäre. Das wird zusammen mit den weiter unten folgenden Eigenschaften der Exponentialfunktion im **alten Skript** bewiesen.

Aufgabe 33.3. Welchen Wert hat $\exp(0)$? Berechnen Sie mit einem Taschenrechner gemäß der Definition in (33.1) einen ungefähren Wert für e .

Eigenschaften der Exponentialfunktion

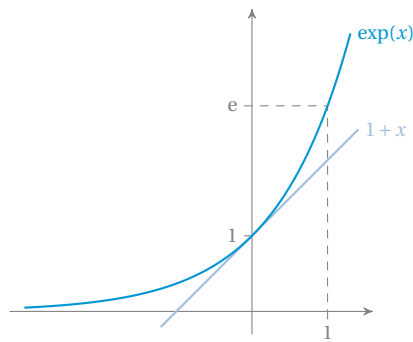
Für alle $w, z \in \mathbb{C}$ gelten die folgenden Aussagen:

- (i) $\exp(w+z) = \exp(w) \exp(z)$
- (ii) $\exp(z) \neq 0$
- (iii) $\exp(z) \in \mathbb{R}^+$ und $\exp(z) \geq 1+z$ für $z \in \mathbb{R}$
- (iv) $\exp(w) < \exp(z)$ für $w, z \in \mathbb{R}$ und $w < z$
- (v) $\exp(qz) = \exp(z)^q$ für $q \in \mathbb{Z}$
- (vi) $\exp(qz) = \exp(z)^q$ für $q \in \mathbb{Q}$ und $z \in \mathbb{R}$

Satz 33.2

„Verkleinert“ man wie in Abschnitt 19 den Definitionsbereich von $\exp$ zu $\mathbb{R}$, so erhält man also eine Funktion von $\mathbb{R}$ nach $\mathbb{R}^+$. Deren Verlauf sieht so aus wie in der folgenden Skizze. Auf der linken Seite bleibt die Kurve immer oberhalb der horizontalen Achse, nach rechts steigt sie sehr stark an.¹⁰

¹⁰Nach Korollar 31.3 steigt $\exp$ stärker als jedes Polynom.



Aufgabe 33.4. Warum gibt es wohl zwei Versionen der Formel $\exp(qz) = \exp(z)^q$ in Satz 33.2?

Definition

Für $z \in \mathbb{C}$ definieren wir e^z als $\exp(z)$.

Wir können also nun beispielsweise sagen, was e^π oder e^i bedeuten sollen, während diese Ausdrücke vorher keinen Sinn ergeben hätten – wohingegen Terme der Form e^q mit rationalen Exponenten q definiert waren. $z \rightarrow e^z$ ist damit eine auf ganz $\mathbb{C}$ definierte Funktion, die nach Satz 33.2 für rationale Argumente mit der bisherigen Definition übereinstimmt.

Aufgabe 33.5. Geben Sie alle Lösungen der Gleichung $(2x - 6)e^{x^2+3} = 0$ an.

Man kann nun den folgenden Zusammenhang beweisen, der zeigt, dass man den Ausdruck e^λ nicht nur über die Reihe (33.1) berechnen kann, sondern auch als Grenzwert einer Folge.

Satz 33.3

Für alle $\lambda \in \mathbb{R}$ ist e^λ der Grenzwert der Folge $\left(1 + \lambda/n\right)^n_{n \in \mathbb{N}}$.

Aufgabe 33.6. Nehmen wir an, es gäbe eine extrem großzügige Bank, die auf Sparguthaben pro Jahr 100 Prozent Zinsen zahlt und zudem verspricht, dass man den entsprechenden Anteil auch erhält, wenn man das Geld früher abhebt – also z. B. 50 Prozent Zinsen nach einem halben Jahr. Sie legen 1000 Euro an und haben die clevere Idee, Einlage plus Zinsen nach einem halben Jahr abzuheben, das gesamte Geld aber sofort wieder für ein weiteres halbes Jahr anzulegen. Haben Sie nach Ablauf des Jahres dann mehr Geld als jemand, der ebenfalls 1000 Euro eingezahlt, aber einfach das Jahresende abgewartet hat?

Aufgabe 33.7. Man könnte die Geschichte aus Aufgabe 33.6 weiterspinnen und z. B. jeweils schon nach vier Monaten abheben und gleich wieder einzahlen. Oder im Rhythmus von zwei Monaten. Oder noch öfter. Stellen Sie eine Formel dafür auf, wie viel Geld man am Ende hat, wenn man das Jahr in n gleiche Teile zerlegt.

Aufgabe 33.8. Berechnen Sie die Grenzwerte der Folgen

$$\left(1 + 1/n\right)^{n+1}_{n \in \mathbb{N}} \quad \text{und} \quad \left(1 + 1/n\right)^{2n}_{n \in \mathbb{N}}.$$

Aufgabe 33.9. Geben Sie den Grenzwert von $\left(\frac{n}{n+1}\right)^n$ an. (Tipp: Betrachten Sie den Kehrwert des Bruchs.)

Aufgabe 33.10. Geben Sie den Grenzwert dieser Folge an:

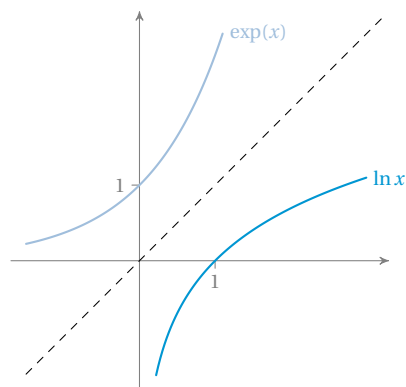
$$\left(\left(\frac{n^2}{n^2 + 2n + 1} \right)^n \right)_{n \in \mathbb{N}}.$$

Nach Punkt (iv) von Satz 33.2 ist die Exponentialfunktion, wenn man sie auf $\mathbb{R}$ einschränkt, injektiv. Daher können wir die **Umkehrfunktion** definieren.

Die auf $\mathbb{R}^+$ definierte Umkehrfunktion der reellen Exponentialfunktion wird als **natürlicher Logarithmus** bezeichnet und $\ln$ geschrieben.¹¹

Definition

Da sich (siehe Seite 164) die Funktionsgraphen von Umkehrfunktionen reeller Funktionen durch Spiegelung an der Hauptdiagonalen ergeben, sieht der Logarithmus folgendermaßen aus.



Als Folgerung aus Satz 33.2 erhält man sofort, dass der Logarithmus gewissermaßen Multiplikation und Division in die einfacheren Operationen Addition und Subtraktion „übersetzen“ kann.

Für alle $x_1, x_2 \in \mathbb{R}^+$ gilt $\ln x_1 x_2 = \ln x_1 + \ln x_2$, sowie $\ln(x_1 / x_2) = \ln x_1 - \ln x_2$.

Korollar 33.4

Ist a eine positive reelle Zahl und $q \in \mathbb{Q}$, dann gilt nach Satz 33.2

$$a^q = \exp(\ln a)^q = \exp(q \ln a).$$

Analog zur Definition von e^z ergibt daher auch die folgende Definition Sinn, stimmt für rationale Exponenten mit der gewohnten überein, liefert nun jedoch (für festes a) eine auf ganz $\mathbb{C}$ definierte Funktion.

¹¹In mathematischen Fachtexten wird statt $\ln$ oft auch $\log$ verwendet. Auch in vielen Programmiersprachen heißt die entsprechende Funktion $\log$.

Definition

Für $a > 0$ und $z \in \mathbb{C}$ definieren wir a^z als $\exp(z \ln a)$ bzw. $e^{z \ln a}$.

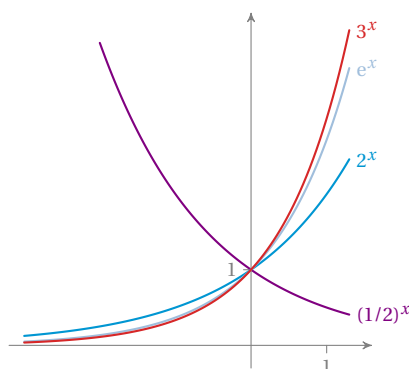
Damit erhalten wir unter anderem den erwünschten Effekt, dass sich beim Rechnen nichts ändert:

Lemma 33.5

Die Rechenregeln aus Lemma 5.8 gelten für beliebige reelle Basen a und b sowie beliebige komplexe Exponenten m und n , sofern die entsprechenden Ausdrücke definiert sind.

Exponentialfunktion

Funktionen der Art $x \mapsto a^x$ werden allgemein als *Exponentialfunktionen* bezeichnet. Einige Beispiele zeigt die folgende Skizze. Beachten Sie, dass alle diese Funktionen sich im Punkt $(0, 1)$ schneiden. Außerdem folgt aus den Rechenregeln in Lemma 33.5, dass der Funktionsgraph von $(1/a)^x$ sich durch Spiegelung des Graphen von a^x an der vertikalen Achse ergeben muss.



Da diese Funktionen offenbar alle injektiv sind, kann man nun auch allgemeine Logarithmen definieren.

Definition

Sei $a > 0$ und $a \neq 1$. Die auf $\mathbb{R}^+$ definierte Umkehrfunktion von $x \mapsto a^x$ wird als **Logarithmus zur Basis a** bezeichnet und $\log_a$ geschrieben.¹²

Da es erfahrungsgemäß oft zu gewissen Verwirrungen kommt, sei hier ein wichtiger Punkt explizit hervorgehoben. Wenn Sie Gleichungen wie $x + y = a$ oder $xy = a$ nach x auflösen, dann erhalten Sie $x = a - y$ bzw. $x = a/y$, weil Subtraktion und Division die Umkehrung von Addition und Multiplikation sind – siehe Seite 41. Und wenn Sie diese Gleichungen nach y auflösen, dann ergeben sich $y = a - x$ bzw. $y = a/x$. Das ist dieselbe Rechnung, bei der nur die Rollen von x und y vertauscht wurden. Und das liegt daran, dass Addition und Multiplikation **kommutativ** sind.

Das Potenzieren ist aber *nicht* kommutativ – z. B. ist 2^3 nicht dasselbe wie 3^2 . Daher ist es bei einer Gleichung wie $x^y = a$ *nicht* egal, ob Sie nach x oder nach y auflösen. Nach x aufgelöst erhält man $x = a^{1/y}$ und beantwortet die Frage, welche Zahl

¹²In technischen Texten findet man auch die Abkürzungen lg für $\log_{10}$ sowie lb oder ld für $\log_2$. Und ln ist natürlich dasselbe wie $\log_e$.

zur y -ten Potenz a ergibt.¹³ Nach y aufgelöst erhält man jedoch $y = \log_x a$ und beantwortet die Frage, mit was man x potenzieren muss, damit a herauskommt.

Aufgabe 33.11. Geben Sie ohne Zuhilfenahme eines Taschenrechners oder anderer Hilfsmittel die Werte $\log_2 64$ und $\log_{10} 100\,000$ an.

Korollar 33.4 gilt für alle Logarithmen. Und man kann noch weitere angenehme Eigenschaften beweisen.

Für alle $a, b, x, y > 0$ mit $a, b \neq 1$ gelten die folgenden Aussagen:

- (i) $\log_a x = \log_b x / \log_b a$
- (ii) $\log_a xy = \log_a x + \log_a y$
- (iii) $\log_a (x/y) = \log_a x - \log_a y$
- (iv) $\log_a x^y = y \log_a x$

Lemma 33.6

Aufgabe 33.12. Viele Taschenrechner kennen nur die zwei Logarithmusfunktionen $\ln$ und $\log_{10}$. Wie würden Sie mit so einem Taschenrechner $\log_7 42$ berechnen?

Aufgabe 33.13. Sei $a = 500^{-\frac{1}{2}}$ und $b = \sqrt{5}/a$. Berechnen Sie $c = \log_{10} b + \log_{10} 20$. (Arbeiten Sie mit Zettel und Bleistift. Das Ergebnis ist eine rationale Zahl und kann daher exakt angegeben werden.)

Aufgabe 33.14. Seien b und c positive reelle Zahlen und $x = e^{b \ln c}$. Geben Sie einen möglichst einfachen Ausdruck für $x^{-2/b}$ an, in dem weder x und e noch Logarithmen vorkommen.

Aufgabe 33.15. Sei $a > 0$. Geben Sie die Lösung x der Gleichung $a^x = 5/a^2$ in Abhängigkeit von a in möglichst einfacher Form an.

Aufgabe 33.16. Überlegen Sie sich ein Beispiel, das zeigt, dass $\log_2 a^b = a \log_2 b$ im Allgemeinen *nicht* gilt. Wie lautet die korrekte Formel?

Bemerkungen

Beachten Sie, dass in diesem Abschnitt zwei weitere wichtige Folgen zu denen aus Satz 30.2 hinzugekommen sind, nämlich die aus Lemma 33.1 und aus Satz 33.3.

Die Fakultät hat eine höhere Priorität als Punkt- und Strichrechnung. $k \cdot n!$ bedeutet also $k \cdot (n!)$. Sie sollten in so einem Fall allerdings den Multiplikationspunkt nicht weglassen, weil man $kn!$ als $(kn)!$ missverstehen könnte.

Für das Setzen bzw. Weglassen von Klammern gelten beim Logarithmus dieselben Konvention wie bei den trigonometrischen Funktionen. Nur bei der Exponentialfunktion lässt man in der Regel (wohl aus historischen Gründen) die Klammern *nicht* weg.

¹³Wenn $y \in \mathbb{N}^+$ gilt, spricht man bekanntlich auch von der y -ten Wurzel.

Wegen der „[Übersetzungseigenschaften](#)“ wurden Logarithmen Anfang des 17. Jahrhunderts ursprünglich als Rechenhilfe eingeführt und behielten diesen Status bis zur allgemeinen Verfügbarkeit elektronischer Taschenrechner bei. Im Video [Was sind Logarithmen?](#) gibt es dazu weiterführende Informationen. Dort werden auch typische Anwendungen von Logarithmen angesprochen.

★ Exponentialfunktionen und Logarithmen im Computer

In PYTHON gibt es für die Exponentiation drei Möglichkeiten, die sich subtil unterscheiden.

```
import math
2**3, pow(2,3), math.pow(2,3)
```

Die dritte Version rechnet immer mit Fließkommazahlen gemäß der Definition auf Seite 258, während die anderen beiden bei ganzzahligen Argumenten auch ganzzahlige Werte zurückgeben. Außerdem können diese auch mit [PYTHONs komplexen Zahlen](#) umgehen. Man beachte die Klammern, die in diesem Fall notwendig sind:

```
(2+5j)**3j, pow(2+5j,3j)
```

Dass pow auch [modulare Arithmetik](#) beherrscht, [wissen wir schon](#).

Die Standardbibliothek MATH beinhaltet die natürliche Exponentialfunktion, Logarithmen und die eulersche Zahl. Ruft man log mit zwei Argumenten auf, dann ist das zweite die Basis.

```
from math import *
exp(1), e, log(e), log(1000,10)
```

Wenn Sie mit komplexen Zahlen arbeiten wollen, müssen Sie [die Bibliothek CMATH](#) verwenden.¹⁴

```
import cmath
x=cmath.exp(2+3j)
x, cmath.log(x)
```

Zusätzlich gibt es Spezialfunktionen wie expm1 oder log10, die in Grenzfällen genauer rechnen. Für Einzelheiten sei auf [die Dokumentation](#) verwiesen.

```
log10(1000), log(1000,10)
```

¹⁴[Komplexe Logarithmen](#) haben wir in diesem Skript gar nicht definiert.

34. Stetige Funktionen

Der Begriff der **Stetigkeit** präzisiert mathematisch in gewissem Sinne eine Grundannahme der Philosophie und Naturwissenschaft seit der Antike, die man als *Natura non facit saltus* – „Die Natur macht keine Sprünge“ – bezeichnet. Heutzutage betrachtet man Stetigkeit als eine **topologische** Eigenschaft, die man unabhängig von der **Differenzierbarkeit** einer Funktion untersuchen kann.

Sei A eine Menge von komplexen Zahlen. Wir sagen, dass (a_n) eine **Folge** in A ist, wenn alle Folgenglieder Elemente von A sind. a sei eine Zahl, für die mindestens eine Folge in A existiert, die gegen a konvergiert. Ferner sei f eine Funktion von A nach $\mathbb{C}$. Gilt $\lim_{n \rightarrow \infty} f(a_n) = L$ für *alle* Folgen (a_n) in $A \setminus \{a\}$ mit $\lim_{n \rightarrow \infty} a_n = a$, dann nennt man L den **Grenzwert** von f an der Stelle a und schreibt $\lim_{x \rightarrow a} f(x) = L$.

Definition

Es ist wichtig, den Unterschied zur Grenzwertdefinition für **Folgen** zu verstehen. Man kann eine Folge zwar als Funktion im Sinne dieser Definition interpretieren, aber ∞ ist im Gegensatz zu a keine Zahl. Es ist zwar möglich, beide Definitionen **unter einen Hut** zu bringen, aber das wäre zu abstrakt für das Niveau dieses Skripts.

Ganz wesentlich ist außerdem, dass a nicht notwendig im **Definitionsbereich** von f liegen muss. Man kann also vom Grenzwert an der Stelle a sprechen, obwohl der Funktionswert $f(a)$ evtl. gar nicht definiert ist. Beispielsweise ist die Funktion

$$f: \begin{cases} \mathbb{R} \setminus \{2\} \rightarrow \mathbb{R} \\ x \mapsto \frac{x^2 - 4}{x - 2} \end{cases}$$

für $x = 2$ nicht definiert. Ist jedoch (a_n) eine Folge in $\text{dom}(f)$, so ist sind natürlich alle Folgenglieder von $(f(a_n))$ definiert. Gilt ferner $\lim_{n \rightarrow \infty} a_n = 2$, so erhalten wir mit der dritten **binomischen Formel**

$$f(a_n) = \frac{a_n^2 - 4}{a_n - 2} = \frac{(a_n - 2)(a_n + 2)}{a_n - 2} = a_n + 2 \xrightarrow{n \rightarrow \infty} 2 + 2 = 4.$$

Also gilt $\lim_{x \rightarrow 2} f(x) = 4$.

Wir werden nicht mit Grenzwerten rechnen. Da es sich aber um einen grundlegenden Begriff der Analysis handelt, sollte man davon zumindest einmal gehört haben. Ohne Grenzwerte kann man außerdem weitere wichtige Konzepte wie z. B. das folgende nicht präzise definieren.

Sei $f: A \rightarrow \mathbb{C}$ mit $A \subseteq \mathbb{C}$. Man sagt, dass f **stetig** an der Stelle $a \in A$ ist, wenn $\lim_{x \rightarrow a} f(x) = f(a)$ gilt. Ist $B \subseteq A$ und f stetig an allen Stellen $a \in B$, dann nennt man f **stetig** auf B . Ist f stetig auf dem gesamten Definitionsbereich A , so nennt man f einfach **stetig**.

Definition

Für Funktionen, die reelle auf reelle Zahlen abbilden, ist eine in den meisten Fällen ausreichende intuitive Vorstellung, dass sie stetig sind, wenn man ihren Funktionsgraphen „in einem Strich“ (also ohne Absetzen) zeichnen kann.

Aufgabe 34.1. Zeichnen Sie eine Skizze des Funktionsgraphen von

$$f: \begin{cases} \mathbb{R} \setminus \{1\} \rightarrow \mathbb{R} \\ x \mapsto \begin{cases} x & x < 1 \\ x+1 & x \geq 1 \end{cases} \end{cases}.$$

Die Funktion f aus der Aufgabe 34.1 ist beispielsweise an der Stelle 1 *nicht* stetig. Dafür kann man etwa die Folge $(1 - 1/n)$ betrachten, die gegen 1 konvergiert. Setzt man die Folgenglieder in f ein, so erhält man wieder die Folge $(1 - 1/n)$. Der Funktionswert von f an der Stelle 1 ist jedoch 2 und nicht 1. (An allen anderen Stellen ist f jedoch stetig.)

Stellen, an denen Funktionen nicht stetig sind, entstehen häufig durch **Fallunterscheidung** wie in obigen Beispiel. Es gibt zwar noch raffiniertere Methoden, Unstetigkeit zu erzwingen, aber das spielt fast nie eine Rolle, wenn man nicht gerade Mathematik studiert. Die gute Nachricht ist, dass die Funktionen, mit denen wir es täglich zu tun haben, alle stetig sind. Das wird im Folgenden beschrieben.

Definition

Sind f und g komplexwertige Funktionen mit demselben Definitionsbereich $A \subseteq \mathbb{C}$, so definiert man die Funktion $f + g$ durch

$$f + g: \begin{cases} A \rightarrow \mathbb{C} \\ z \mapsto f(z) + g(z) \end{cases}.$$

Analog werden $f - g$ und $f \cdot g$ (bzw. fg) definiert. Gilt $g(z) \neq 0$ für alle $z \in A$, so kann man ebenso f/g definieren. Ist schließlich f die konstante Funktion $f(x) = \lambda$, so schreibt man z. B. statt fg auch λg .

Definition

Sei n eine natürliche Zahl und seien a_i bis a_n komplexe Zahlen mit $a_n \neq 0$. Eine Funktion der Form

$$: \begin{cases} \mathbb{C} \rightarrow \mathbb{C} \\ z \mapsto a_0 + a_1 z + a_2 z^2 + \cdots + a_n z^n \end{cases}$$

wird **Polynom** vom **Grad** n genannt. Die a_i nennt man die **Koeffizienten** des Polynoms. Ergänzend wird die Funktion von $\mathbb{C}$ nach $\mathbb{C}$, die z konstant auf 0 abbildet, als **Nullpolynom** bezeichnet und bekommt den Grad $-\infty$ zugeordnet.

Aufgabe 34.2. Machen Sie sich klar, dass Summen, Differenzen und Produkte von Polynomen wieder Polynome sind. Welchen Grad hat das Polynom $p \cdot q$, wenn p und q den Grad m bzw. n haben?

Aufgabe 34.3. Geben Sie Polynome p und q vom Grad 2 an, deren Summe $p + q$ den Grad 1 hat.

Alle Polynome, alle **Exponentialfunktionen** und **Logarithmen**, die **Wurzelfunktionen**, die **trigonometrischen Funktionen** und die **Arkusfunktionen** sind stetig. Und sind f und g stetig, so sind auch die Funktionen $f + g$, $f - g$, $f \cdot g$, f/g und $f \circ g$ stetig.

Satz 34.1

Man kann sich also aus einem „Baukasten“ bekannter Funktionen unendlich viele andere zusammenbasteln und erhält immer wieder stetige Funktionen. Hier drei Beispiele:

$$\begin{aligned} &: \begin{cases} \mathbb{R} \rightarrow \mathbb{R} \\ x \mapsto \frac{3 \sin x}{x^2 + 1} \end{cases} &: \begin{cases} \mathbb{C} \rightarrow \mathbb{C} \\ z \mapsto \exp(z^3 - 2z) \end{cases} &: \begin{cases} \mathbb{R}^+ \setminus \{1\} \rightarrow \mathbb{R} \\ x \mapsto \frac{1}{1-x} \cdot \ln x \end{cases} \end{aligned}$$

Alle so entstandenen Funktionen nennt man typischerweise **elementar**; allerdings wird der Begriff in der Literatur nicht eindeutig gehandhabt. Wichtig ist jedoch, dass immer der Definitionsbereich zu beachten ist. Der Tangens kann beispielsweise an der Stelle $\pi/2$ nicht stetig sein und die Funktion $x \mapsto 1/x$ nicht an der Stelle 0.

elementare Funktion

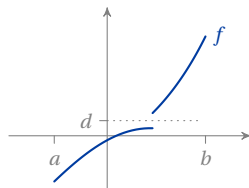
Die intuitive Vorstellung, dass stetige Funktionen keine „**Sprünge**“ machen, dass sie also keine Werte „auslassen“, kann man (zumindest im Reellen) folgendermaßen präzisieren.

Zwischenwertsatz

Ist $f: [a, b] \rightarrow \mathbb{R}$ stetig (wobei $a < b$ gelten soll), so gibt es für alle Zahlen d zwischen $f(a)$ und $f(b)$ ein $c \in [a, b]$ mit $f(c) = d$.

Satz 34.2

Die folgende Skizze zeigt, wie eine Funktion f eine Zahl d „überspringt“, sie also nicht als Funktionswert annimmt, obwohl d zwischen $f(a)$ und $f(b)$ liegt. Bei stetigen Funktionen kann so etwas nicht passieren.



Aufgabe 34.4. Von einer Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ weiß man nur, dass sie stetig ist und dass $f(1/k) = (-2)^k$ für alle $k \in [5]$ gilt. Wie viele Nullstellen muss f mindestens haben?

Aufgabe 34.5. Seien f und g stetige Funktionen von $\mathbb{R}$ nach $\mathbb{R}$. Ferner seien $f(3)$, $g(3)$ und $f(4)$ positiv. Aus einer der folgenden sechs Aussagen kann man folgern, dass die Funktion $h = f \cdot g$ im Intervall $(3, 4)$ mindestens eine Nullstelle hat. Welche ist es und warum?

$$g(4) > 0 \quad g(4) \geq 0 \quad g(4) < 0 \quad g(4) \leq 0 \quad g(4) = 0 \quad g(4) = f(4)$$

Der **Zwischenwertsatz** sichert für viele Gleichungen die Existenz einer Lösung und liefert gleichzeitig ein Verfahren, um Näherungswerte für so eine Lösung zu bestimmen, das man **Bisektion** nennt:

Bisektion

Nullstelle

- Jede Gleichung mit einer Unbekannten x lässt sich so umformen, dass auf einer Seite der Gleichung 0 steht. Interpretiert man die andere Seite als Funktion f von x , so bekommt die Gleichung die Form $f(x) = 0$. Beim Lösen der Gleichung geht es also darum, eine sogenannte **Nullstelle** von f zu finden.
- Ist f auf einem Intervall $[a_0, b_0]$ definiert und stetig und haben $f(a_0)$ und $f(b_0)$ unterschiedliche Vorzeichen, so gibt es nach dem Zwischenwertsatz eine Zahl $c \in (a_0, b_0)$ mit $f(c) = 0$. Damit ist klar, dass eine Lösung existiert.
- Nun setzt man $c_0 = (a_0 + b_0)/2$ und berechnet $f(c_0)$. Gilt $f(c_0) = 0$, so ist man fertig. Das ist allerdings sehr unwahrscheinlich. Anderenfalls haben aber entweder a_0 und c_0 oder c_0 und b_0 unterschiedliche Vorzeichen. Man kann dann $[a_0, b_0]$ durch $[a_0, c_0]$ oder $[c_0, b_0]$ ersetzen, dieses neue Intervall $[a_1, b_1]$ nennen und den Vorgang wiederholen, indem man die Mitte c_1 von $[a_1, b_1]$ berechnet. Und so weiter.
- Damit erhält man eine Folge (c_n) , die offenbar gegen eine Nullstelle von f konvergiert. In der Praxis wird man nach endlich vielen Schritten aufhören, wenn der Abstand zwischen den Intervallgrenzen so gering ist, dass man mit der Genauigkeit der Näherung zufrieden ist.

Für die Gleichung $x^5 + 1 = x$, für die es übrigens **keine Lösungsformel** gibt, wäre f die durch $f(x) = x^5 - x + 1$ definierte Funktion. Es gilt $f(-2) = -29$ und $f(2) = 31$. Mit den Startwerten -2 und 2 würden die ersten Schritte so aussehen:

k	a_k	b_k
0	-2	2
1	-2	0
2	-2	-1
3	-1.5	-1
4	-1.25	-1
5	-1.25	-1.125
6	-1.1875	-1.125
7	-1.1875	-1.15625
8	-1.171875	-1.15625
9	-1.171875	-1.1640625
10	-1.16796875	-1.1640625
11	-1.16796875	-1.166015625
12	-1.16796875	-1.1669921875
13	-1.16748046875	-1.1669921875
14	-1.16748046875	-1.167236328125
15	-1.1673583984375	-1.167236328125

Eine Lösung der Gleichung¹⁵ liegt also zwischen a_{15} und b_{15} .

¹⁵Es gibt insgesamt fünf, aber die anderen vier sind echt komplex.

Aufgabe 34.6. Greifen Sie sich einen Taschenrechner und rechnen Sie zumindest die ersten Zeilen der obigen Tabelle nach, um sicherzustellen, dass Sie das Verfahren verstanden haben.

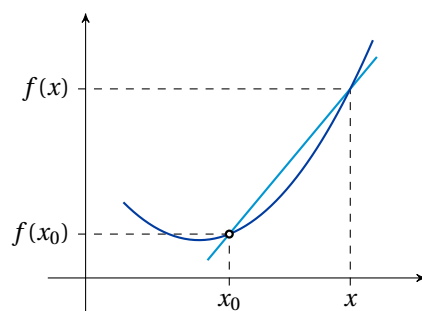
Aufgabe 34.7. Wenn man wie in diesem Beispiel bis zum 15. Schritt gerechnet hat, wie viele **Nachkommastellen** der Lösung sollte man dann angeben?

Aufgabe 34.8. Berechnen Sie mit dem Bisektionsverfahren eine Näherung für die reelle Lösung der Gleichung $x^5 + 2x^2 = x + 8$ auf zwei Stellen nach dem Komma.

★Aufgabe 34.9. Programmieren Sie das Bisektionsverfahren in einer Programmiersprache Ihrer Wahl.

35. Differenzieren

Die Idee der Ableitung ist in Kombination mit der des **Integrals** die wohl wichtigste Innovation der letzten tausend Jahre in der Mathematik. Motivieren kann man diese Idee dadurch, dass man an den Funktionsgraphen einer Funktion f im Punkt $(x_0, f(x_0))$ eine Tangente anlegen will, deren Steigung man ermitteln möchte.¹⁶



Die in der Skizze eingezeichnete Gerade hat die Steigung $(f(x) - f(x_0)) / (x - x_0)$. Sie ist *nicht* die Tangente, aber es sieht so aus, als würde man nach und nach bessere Näherungswerte für die Steigung der Tangente erhalten, wenn x näher an x_0 heranrückt.

Sei $f: A \rightarrow \mathbb{R}$ mit $A \subseteq \mathbb{R}$ und x_0 ein Element von A mit der Eigenschaft, dass für eine Zahl $\varepsilon > 0$ das Intervall $(x_0 - \varepsilon, x_0 + \varepsilon)$ eine Teilmenge von A ist. f wird **differenzierbar** an der Stelle x_0 genannt, wenn der Grenzwert

$$\lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0} \quad (35.1)$$

(als *Zahl!*) existiert. Für diesen Grenzwert schreibt man dann

$$\frac{df}{dx}(x_0) \quad \text{oder} \quad \frac{d}{dx}f(x_0) \quad \text{oder} \quad \left. \frac{df(x)}{dx} \right|_{x=x_0}$$

Definition

¹⁶In diesem Abschnitt und dem folgenden werden wir uns nur mit Funktionen beschäftigen, die reelle auf reelle Zahlen abbilden. Man kann viele der angesprochenen Konzepte **auch auf komplexe Funktionen übertragen**, aber das würde uns zu weit führen und die grafische Intuition erschweren.

und nennt ihn die **Ableitung** bzw. den **Differentialquotienten** von f an der Stelle x_0 . Ist klar, welche Variable gemeint ist, so verwendet man für die Ableitung auch die Kurzschreibweise $f'(x_0)$. Ist B eine Teilmenge von A und f differenzierbar in allen Punkten $x_0 \in B$, dann sagt man, dass f **differenzierbar auf B** ist. Handelt es sich bei B um den Definitionsbereich A von f , so sagt man einfach, dass f **differenzierbar** ist.

Zu beachten ist der erste Satz, in dem gefordert wird, dass x_0 nicht nur ein Element von A ist, sondern dass gewissermaßen „genügend Nachbarschaft“ von x_0 ebenfalls zu A gehört.¹⁷ Anderenfalls ergibt das ganze Konzept keinen Sinn.

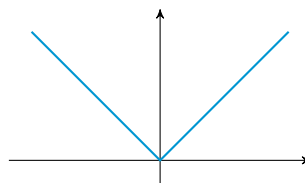
Der geniale „Trick“ bei dieser Definition ist, dass die Funktion in (35.1) an der Stelle $x = x_0$ nicht definiert ist (weil sonst durch null dividiert würde).

Durch simples Einsetzen in die Definition kann man sich überzeugen, dass konstante Funktionen überall differenzierbar mit der Ableitung 0 sind. Ebenso ist die Identität überall differenzierbar mit der Ableitung 1.

Aufgabe 35.1. Rechnen Sie anhand der Definition nach, dass die obigen Aussagen über konstante Funktionen und die Identität korrekt sind.

Aufgabe 35.2. Berechnen Sie nur mithilfe der obigen Definition (also ohne die Verwendung Ihres Vorwissens aus der Schule) die Ableitung der Funktion $x \mapsto x^2$ an der Stelle x_0 . (Hinweis: Dritte **binomische Formel**.)

Ein Beispiel für eine nicht überall differenzierbare Funktion ist die Betragsfunktion $x \mapsto |x|$. An der Stelle $x_0 = 0$ ist diese Funktion nicht differenzierbar, denn der Quotient $(|x| - |0|)/(x - 0)$ ist positiv für $x > 0$ und negativ für $x < 0$. Darum würde man mit den Folgen $(1/n)$ und $(-1/n)$, die beide gegen 0 konvergieren, unterschiedliche Grenzwerte erhalten, wenn man sie in diesen Quotienten einsetzen würde. Man „sieht“ ja auch, dass es an dieser Stelle keine eindeutige Tangente gibt.



Ist f an der Stelle x_0 differenzierbar, so kann man das auf verschiedene Arten interpretieren:

- Die Ableitung an der Stelle x_0 ist die Steigung der Tangente an den Funktionsgraphen von f durch den Punkt $(x_0, f(x_0))$.
- Wenn die x -Achse einen Zeitverlauf darstellt, dann ist $f'(x_0)$ die *momentane Änderungsrate* der Funktionswerte zum Zeitpunkt x_0 . (Ist $f(x)$ beispielsweise die bis zur Zeit x zurückgelegte Strecke, dann ist $f'(x_0)$ die Geschwindigkeit zum Zeitpunkt x_0 – also die momentane Änderungsrate der Strecke.)

¹⁷Der Fachausdruck ist, dass x_0 ein *innerer Punkt* von A ist.

- Differenzierbare Funktionen sehen „unter dem Mikroskop“ wie Geraden aus. Grafisch wird dies in dem Video [Die Ableitung unter dem Mikroskop](#) demonstriert.

Um Schreibarbeit zu sparen, sollen im Folgenden die Funktionen, um die es geht, immer solche sein, die reelle auf reelle Zahlen abbilden. Bei den Stellen, an denen differenziert wird, soll es zudem natürlich immer um Punkte gehen, auf die die obige Definition zutrifft.

Die Betragsfunktion hat gezeigt, dass aus Stetigkeit nicht notwendig Differenzierbarkeit folgt.¹⁸ Umgekehrt gilt das aber schon:

Ist f an der Stelle x_0 differenzierbar, dann ist f dort auch stetig.

Lemma 35.1

Es folgen nun diverse Rechenregeln für die Ableitung, die hoffentlich größtenteils schon aus der Schule bekannt sind. Für Beweise sei wie üblich auf das [alte Skript](#) verwiesen.

Sind f und g an der Stelle x_0 differenzierbar und ist λ eine Konstante, so sind auch $f \pm g$ und λf an dieser Stelle differenzierbar und für die Ableitungen gilt $(f \pm g)'(x_0) = f'(x_0) \pm g'(x_0)$ und $(\lambda f)'(x_0) = \lambda f'(x_0)$.

Lemma 35.2

Das ist die [Linearität](#) des Ableitens. Damit erhält man beispielsweise, dass die Funktion $x \mapsto 3x + 5$ überall differenzierbar mit der Ableitung 3 ist.

Leibniz- bzw. Produktregel

Sind f und g an der Stelle x_0 differenzierbar, so ist auch ihr Produkt fg an dieser Stelle differenzierbar und es gilt $(fg)'(x_0) = f'(x_0)g(x_0) + f(x_0)g'(x_0)$.

Satz 35.3

Aufgabe 35.3. Berechnen Sie die Ableitung von $x \mapsto x^2$ erneut (siehe Aufgabe 35.2), indem Sie diesmal x^2 als Produkt $x \cdot x$ interpretieren und die [Leibnizregel](#) anwenden.

Aufgabe 35.4. Berechnen Sie die Ableitung von x^3 , indem Sie x^3 als $x^2 \cdot x$ darstellen. Berechnen Sie anschließend die Ableitungen von x^4 , x^5 und so weiter, bis Sie überzeugt sind, dass für alle $n \in \mathbb{N}$ die Ableitung von $x \mapsto x^n$ an der Stelle x_0 den Wert nx_0^{n-1} hat.

Aufgabe 35.5. Mit Aufgabe 35.4 und Lemma 35.2 können wir nun jedes [Polynom](#) differenzieren. Wie sieht z. B. die Ableitung von $x \mapsto 3x^3 - 5x^2 + 8x - 2$ an der Stelle x_0 aus?

Ist die Funktion f auf der Menge A differenzierbar, so bezeichnet man die Funktion f' , die jedem $x_0 \in A$ die Ableitung $f'(x_0)$ zuordnet, als **Ableitungs-**

Definition

¹⁸Man kann sogar mit entsprechendem Aufwand Funktionen konstruieren, die *überall* stetig und *nirgends* differenzierbar sind. Beispiele dafür finden Sie in dem Video [Die Monster der Analysis](#).

funktion oder manchmal auch einfach als **Ableitung** von f . Ist f' an der Stelle $x_0 \in A$ ebenfalls differenzierbar, so sagt man, dass f dort **zweimal differenzierbar** ist, und schreibt für diese sogenannte *zweite Ableitung*

$$f''(x_0) \quad \text{oder} \quad \frac{d^2 f}{dx^2}(x_0) \quad \text{oder} \quad \frac{d^2}{dx^2} f(x_0) \quad \text{oder} \quad \left. \frac{d^2 f(x)}{dx^2} \right|_{x=x_0}.$$

Entsprechend werden auch n -te Ableitungen definiert, für die man entweder $f^{(n)}$ schreibt oder eine der obigen Schreibweisen (mit n statt 2) verwendet. Ist eine Funktion n -mal differenzierbar für *alle* $n \in \mathbb{N}$, so spricht man auch von einer **glatten** Funktion.

Aufgabe 35.6. Geben Sie für das Polynom aus Aufgabe 35.5 die ersten vier Ableitungsfunktionen an. Wie sieht die 42. Ableitung aus?

Aufgabe 35.7. Berechnen Sie für die Funktion $f(x) = (x^2 + 3)(x^3 - x)$ die Ableitungsfunktion auf zwei Arten: einmal durch Anwendung der **Leibnizregel** auf das Produkt der beiden Terme und einmal, indem Sie ausmultiplizieren und das resultierende Polynom ableiten. Kommt in beiden Fällen dasselbe heraus?

Satz 35.4

Kehrwertregel

Ist f an der Stelle x_0 differenzierbar und gilt $f(x_0) \neq 0$, so ist auch der Kehrwert $1/f$ an dieser Stelle differenzierbar und es gilt $(1/f)'(x_0) = -f'(x_0)/f(x_0)^2$.

Aufgabe 35.8. Berechnen Sie mit der **Kehrwertregel** die Ableitung von $f(x) = x^{-3}$.

★**Aufgabe 35.9.** Begründen Sie, dass die Regel aus Aufgabe 35.4 auch dann noch gilt, wenn n eine beliebige *ganze* Zahl ist.

Korollar 35.5

Quotientenregel

Sind f und g an der Stelle x_0 differenzierbar und gilt $g(x_0) \neq 0$, so ist auch der Quotient f/g an dieser Stelle differenzierbar und es gilt

$$(f/g)'(x_0) = \frac{f'(x_0)g(x_0) - f(x_0)g'(x_0)}{g(x_0)^2}.$$

Aufgabe 35.10. Berechnen Sie die Ableitung von $f(x) = (x^2 + 2x + 2)/(3x^2 - 2)$.

Aufgabe 35.11. Wenn Sie sich selbst ähnliche Aufgaben stellen und Ihre Lösungen überprüfen wollen, können Sie z. B. in **Wolfram Alpha** so etwas wie

derivative of $(x^2+2x+2)/(3x^2-2)$

eingeben. Wenn das System Terme nicht ausmultipliziert (wie in diesem Fall den Nenner), dann können Sie es mit einer Eingabe wie

expand $(3x^2-2)^2$

dazu „zwingen“.

Die folgende Aussage ist vielen durch die Merkregel „äußere mal innere Ableitung“ bekannt.

Kettenregel

Sei g an der Stelle x_0 differenzierbar und f an der Stelle $g(x_0)$. Dann ist $f \circ g$ an der Stelle x_0 differenzierbar und es gilt $(f \circ g)'(x_0) = f'(g(x_0))g'(x_0)$.

Satz 35.6

Aufgabe 35.12. Berechnen Sie die Ableitung von $h(x) = (3x^2 - 5x + 2)^4$ mithilfe der Kettenregel. Sie sollen dabei weder $h(x)$ noch $h'(x)$ ausmultiplizieren.

Die bisherigen Ableitungsregeln werden im Folgenden noch einmal im „Spickzettelformat“ zusammengefasst.

Linearität des Differenzierens	$(\lambda f)' = \lambda f'$
	$(f + g)' = f' + g'$
Produktregel	$(fg)' = f'g + fg'$
Ableitung des Kehrwerts	$(1/f)' = -f'/f^2$
Quotientenregel	$(f/g)' = (f'g - fg')/g^2$
Kettenregel	$(f \circ g)' = g'(f' \circ g)$

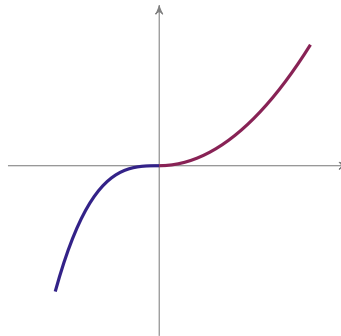
Im gleichen knappen Format folgen nun die Ableitungen der wichtigsten Funktionen. x ist immer die Variable, nach der differenziert wird, a und r sind Konstanten. Wir setzen natürlich voraus, dass die Funktionen sinnvoll definiert sind, dass also etwa bei a^x die Konstante a nicht negativ ist. Die ausgegrauten Zeilen sind Spezialfälle anderer Zeilen und müssten eigentlich nicht aufgeführt werden.

Funktion	Ableitung
a	0
x	1
x^r	rx^{r-1}
e^x	e^x
a^x	$a^x \cdot \ln a$
$\ln x$	x^{-1}
$\log_a x$	$(x \ln a)^{-1}$
$\sin x$	$\cos x$
$\cos x$	$-\sin x$
$\tan x$	$1 + \tan^2 x = (\cos^2 x)^{-1}$
$\arcsin x$	$(\sqrt{1-x^2})^{-1}$
$\arccos x$	$-(\sqrt{1-x^2})^{-1}$
$\arctan x$	$(1+x^2)^{-1}$

Mit den Regeln aus diesem Abschnitt wissen Sie über das Ableiten **elementarer Funktionen** mehr oder weniger alles, was man wissen muss. Und mithilfe der besagten Regeln wird Differenzieren (leider) zu einem mechanischen Vorgang,

den man ohne Nachdenken durchführen kann. Mehr dazu im Video [Wie man Computern das Ableiten beibringt](#).

Gibt es denn überhaupt Funktionen, die wir mit den hier vorgestellten Methoden *nicht* differenzieren können? Die gibt es in der Tat. Einmal gibt es Funktionen, die nicht elementar sind. Aber auch dann, wenn eine Funktion durch **Fallunterscheidung** aus elementaren Funktionen *zusammengesetzt* wurde, versagen unsere Mittel. Die unten skizzierte Funktion ist z. B. links von null als x^3 definiert, rechts davon als $x^2/2$. Im Punkt $x = 0$ hat sie die Ableitung 0, aber unsere Rechenregeln von oben würden uns beim Ermitteln dieses Wertes nicht helfen.



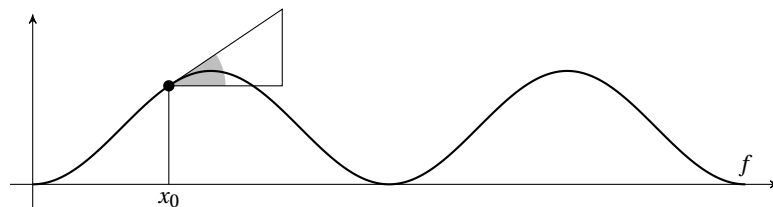
Aufgabe 35.13. Berechnen Sie die Ableitungen der beiden Funktionen $x \mapsto \exp(x^2)$ und $x \mapsto \exp(x)^2$.

Aufgabe 35.14. Die **Exponentialfunktion** hat die bemerkenswerte Eigenschaft, ihre eigene Ableitung zu sein. Durch **Komposition** mit $\exp$ kann man weitere Funktionen konstruieren, die ihre eigene Ableitung sind. Fallen Ihnen Beispiele ein?

Aufgabe 35.15. Berechnen Sie die Ableitungen von $\ln(x^2 + 1)$ und $(\ln x)^2 + 1$.

Aufgabe 35.16. Geben Sie die Ableitungen von $\sin(x^2 + 1)$ sowie von $\ln \sin x$ an.

Aufgabe 35.17. Die Funktion f ist durch $f(x) = \sin^2 x$ definiert. Geben Sie die Ableitung sowie den (grau eingezeichneten) Steigungswinkel φ des Funktionsgraphen von f an der Stelle $x_0 = 6/5$ jeweils auf zwei Stellen nach dem Komma an.



Aufgabe 35.18. Geben Sie die Menge aller Nullstellen der Ableitungsfunktion von $x \mapsto (x^2 + x - 12)^6$ an.

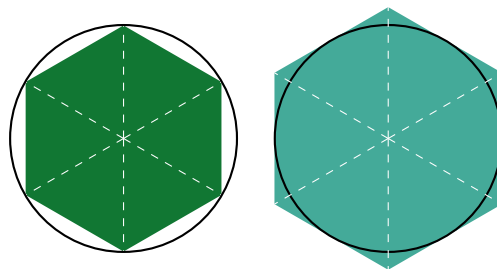
Bemerkungen

Beachten Sie, dass $f'(x_0)$ eine Zahl ist, wenn x_0 eine Zahl ist. Ein typisches Missverständnis ist, dass beispielsweise nach der Ableitung der Funktion $x \mapsto x^2$ an der Stelle 3 gefragt wird und man die Antwort „ $2x$ “ bekommt. Das ist falsch! Die Antwort ist 6.

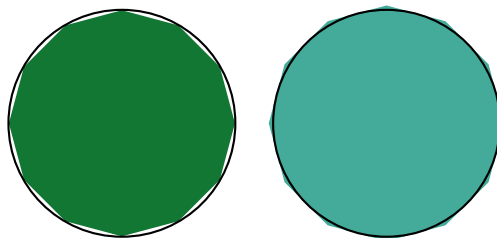
36. Integrieren

Während die Idee des Differenzierens erst in der Neuzeit in Europa entstanden ist, gibt es für das Integrieren Vorläufer, die man bis in die Antike zurückverfolgen kann. (Ausführlich wird das z. B. in dem Video [Das Geheimnis des Archimedes](#) erläutert.) Man kann Integrieren nämlich als das Berechnen von Flächeninhalten interpretieren.

Schon vor der Blütezeit der griechischen Mathematik war bekannt, dass man die Flächen von **Polygonen** prinzipiell immer berechnen kann. Wie ist es aber mit Figuren, deren Seiten nicht alle gerade sind? Die einfachste Figur, die kein Polygon ist, ist der Kreis. **Archimedes** hatte die folgende Idee, seine Fläche zu berechnen: Man lege ein regelmäßiges Sechseck so in den Kreis hinein, dass seine Ecken alle auf dem Kreis liegen. Ein anderes regelmäßiges Sechseck wird so konstruiert, dass seine sechs Seiten den Kreis gerade so von außen berühren.



Beide Flächen lassen sich verhältnismäßig einfach berechnen; und offenbar ist die erste Fläche etwas kleiner als die des Kreises, die zweite etwas größer. Damit hat man den Wert, den man eigentlich haben will, schon mal eingegrenzt. Nun macht man dasselbe noch einmal, aber mit mehr Ecken:



Man bekommt wieder zwei Werte und es ist schon von der Zeichnung her offensichtlich, dass man den tatsächlichen Wert der Kreisfläche jetzt noch genauer eingrenzen kann. Zudem ist auch klar, wie man, wenn man den Aufwand nicht scheut, noch bessere Werte bekommt: man erhöht die Anzahl der Ecken erneut

und erhält so nach und nach immer bessere Näherungswerte für die Fläche des Kreises. Die Grundidee hinter diesem Verfahren hat man bis heute beibehalten:

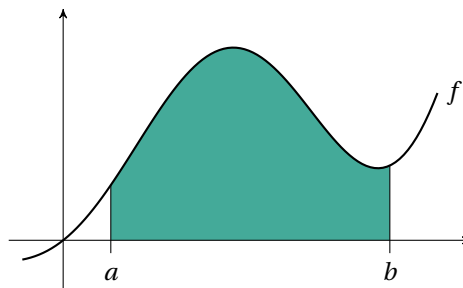
- Man approximiert beliebige geometrische Figuren durch *Polygone*.
- Man verfeinert die so erhaltene Schätzung, indem man die Polygone weiter *unterteilt*.

Der „moderne“ Ansatz sieht allerdings folgendermaßen aus:

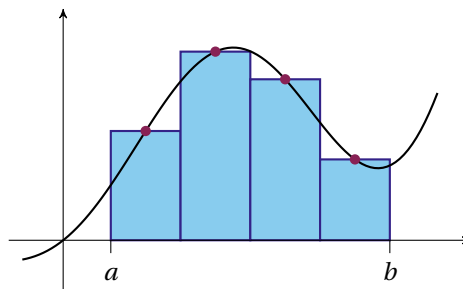
- Man verwendet die Idee der analytischen Geometrie, geometrische Objekte durch **Koordinaten** zu repräsentieren.
- Man stellt die Begrenzungslinien der Figuren, deren Flächen man berechnen will, durch **Funktionsgraphen** dar.
- Man systematisiert das Verfahren, indem man *Rechtecke* statt irgendwelcher Polygone verwendet.
- Und schließlich, das ist der entscheidende Schritt, geht man bei der schrittweisen Verfeinerung zu **Grenzwerten** über, wodurch man von Approximationen zu *exakten* Werten kommt.

Dieser Ansatz wurde im Prinzip seit den Zeiten von **Newton** und **Leibniz** benutzt, aber erst deutlich später präzisiert, unter anderem 1854 durch **Bernhard Riemann**. Das nach ihm benannte *Riemann-Integral* werden wir uns nun anschauen.

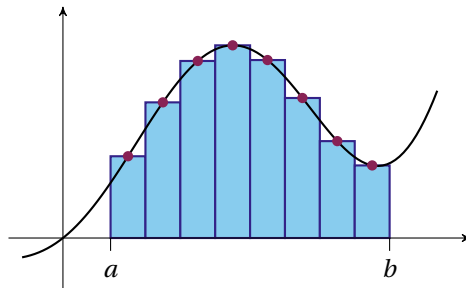
Wir möchten die Fläche unter dem Graphen einer reellwertigen Funktion f zwischen zwei Zahlen a und b berechnen.



Wir zerlegen das Intervall $[a, b]$ in n gleich große Teile und berechnen für jeden Teil den Funktionswert in der Mitte. Daraus ergeben sich n Rechtecke, deren Flächen wir addieren. Für $n = 4$ sieht das so aus:



Die Gesamtfläche der Rechtecke ist eine Näherung für die gesuchte Fläche unter der Kurve. Und diese Näherung wird augenscheinlich besser, wenn wir eine feinere Zerlegung wählen, z. B. $n = 8$:



Formal betrachtet haben wir auf diese Art eine Folge definiert, die jedem n eine Zahl, nämlich die Gesamtfläche der n Rechtecke, zuteilt. Es liegt nahe, zu hoffen, dass diese Folge konvergiert, und zwar gegen die Fläche, die wir ermitteln wollen.

Aufgabe 36.1. In manchen Fällen kann man sich den Aufwand mit der Folge sparen und sieht direkt, welche Fläche sich ergeben muss. Was muss z. B. herauskommen, wenn man die Fläche unter dem Funktionsgraphen der Funktion $x \mapsto 3$ von $x = a$ bis $x = b$ betrachtet? Und was für eine Fläche bekommt man für die Funktion, die x auf $x - a$ abbildet? (Machen Sie sich in beiden Fällen eine Skizze!)

In der (mathematischen) Realität ist die Sache noch etwas komplizierter, weil man Funktionen konstruieren kann, bei denen dieses Verfahren ein offensichtlich falsches Ergebnis liefert. Ein Extrembeispiel ist die sogenannte *Dirichlet-Funktion*, die durch

$$x \mapsto \begin{cases} 1 & x \in \mathbb{Q} \\ 0 & x \in \mathbb{R} \setminus \mathbb{Q} \end{cases}$$

definiert ist. In den Anwendungen werden Sie solchen Ungetümen jedoch selten begegnen. Für viele Funktionen kann eine Fläche wie die im obigen Beispiel im Prinzip mithilfe eines Computerprogramms beliebig genau berechnet werden, wenn man das Intervall $[a, b]$ in genügend kleine Teile zerlegt. Die folgende Definition ist mathematisch nicht präzise genug, für die Praxis ist sie aber in der Regel ausreichend.

Seien a und b reelle Zahlen mit $a < b$ und sei $[a, b]$ eine Teilmenge des Definitionsbereichs der Funktion f . Wenn der oben angedeutete Grenzwert der Rechtecksflächen existiert, so sagt man, dass f auf $[a, b]$ **integrierbar** ist und nennt den Grenzwert das **Integral** von f über $[a, b]$. Die Schreibweise für diese Zahl ist

$$\int_a^b f(x) dx.$$

Definition

Wir setzen außerdem

$$\int_a^a f(x) dx = 0 \quad \text{und} \quad \int_b^a f(x) dx = - \int_a^b f(x) dx.$$

Integrationsvariable

Integrationsgrenzen

Integrand

Das x ist dabei eine austauschbare Variable ähnlich den Laufvariablen in [Summen](#). Man spricht auch von der *Integrationsvariable*. Dass das x hinter dem d noch einmal auftaucht, zeigt, *welche* Variable von a bis b läuft, falls in dem Ausdruck $f(x)$ mehr als eine vorkommt. Die Zahlen a und b sind die *Integrationsgrenzen* und $f(x)$ nennt man in diesem Zusammenhang *Integrand*.

In Aufgabe [36.1](#) haben wir also

$$\int_a^b 3 dx = 3(b-a) \quad \text{und} \quad \int_a^b (x-a) dx = (b-a)^2/2$$

berechnet.

Aufgabe 36.2. Am [Ende des Abschnitts](#) finden Sie eine einfache PYTHON-Funktion, die den Prozess des Approximierens der Fläche durch Rechtecke implementiert. Es ist sicher hilfreich, wenn Sie diese verstehen und ein bisschen mit ihr herumspielen.

Aufgabe 36.3. Überlegen Sie sich, welchen Wert man mittels der beschriebenen Methode (Unterteilung in immer kleinere Rechtecke) für $\int_0^1 g(x) dx$ erhielte, wenn g durch $g(x) = -2$ definiert wäre.

Aufgabe 36.4. Lesen Sie sich die Lösung von Aufgabe [36.3](#) durch und überlegen Sie dann, welchen Wert $\int_{-1}^1 x dx$ haben muss. (Skizze!)

Das von Leibniz eingeführte Zeichen $\int$ für das Integral ist übrigens die kursive Variante des heute nicht mehr gebräuchlichen "[langen s](#)" als Abkürzung des lateinischen Wortes *summa*. Das Zeichen d , auch von Leibniz, ist der erste Buchstabe des lateinischen *differentia*.

Wieso *summa*? Unterteilt man $[a, b]$ in n gleich große Teile und schreibt man Δx für $(b-a)/n$, so ist der Näherungswert, den man durch das oben beschriebene Verfahren erhält, die Summe

$$\sum_{k=1}^n f(x_k) \Delta x,$$

wobei die k Werte $x_k = a - \Delta x/2 + k\Delta x$ über das Intervall $[a, b]$ verteilt sind und Δx die Breite eines Rechtecks ist. Beim Übergang zum Integral wird Δx zum „unendlich kleinen“ dx und x durchläuft *alle* Werte von a bis b . Das Integral ist in diesem Sinne eine *kontinuierliche Summe*.¹⁹

¹⁹Falls Ihnen die Vorstellung einer „unendlich kleinen“ Zahl nicht ganz geheuer ist, ist das nur zu verständlich. Schauen Sie sich in diesem Fall das Video [Was ist Nichtstandardanalysis?](#) an.

Seien a und b reelle Zahlen mit $a < b$. Ist die Funktion f auf $[a, b]$ **stetig**, dann ist f auf $[a, b]$ integrierbar.

Satz 36.1

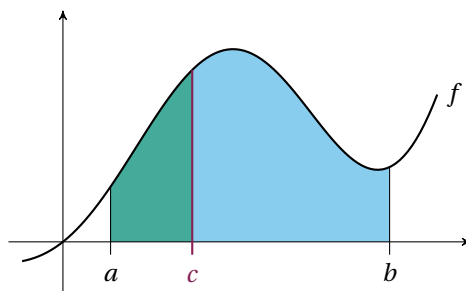
Durch diesen Satz haben wir schon einmal einen großen Vorrat an integrierbaren Funktionen. Wichtig sind allerdings die eckigen Klammern. Die Funktion $1/x$ ist beispielsweise auf dem Intervall $(0, 1]$ stetig, aber *nicht* integrierbar. Weitere integrierbare Funktionen erhält man durch die folgende Eigenschaft, die zumindest von der Anschauung her offensichtlich sein sollte:

Sind a , b und c reelle Zahlen mit $a < c < b$, dann ist eine Funktion f genau dann auf $[a, b]$ integrierbar, wenn sie auf $[a, c]$ und $[c, b]$ integrierbar ist. Es gilt in diesem Fall

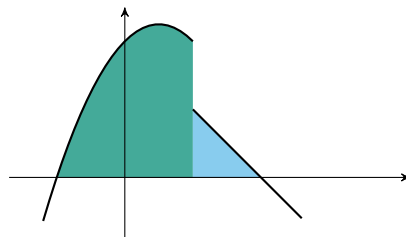
Lemma 36.2

$$\int_a^b f(x) dx = \int_a^c f(x) dx + \int_c^b f(x) dx \quad (36.1)$$

Man kann also direkt nebeneinander liegende Flächenstücke „zusammenkleben“.



Damit sind aber auch Funktionen integrierbar, die nur an wenigen Stellen nicht stetig sind.



(Tatsächlich können es sogar unendlich viele Stellen sein. Aber das behandeln wir hier nicht.)

Aufgabe 36.5. Überlegen Sie sich, dass (36.1) auch dann noch stimmt, wenn man die Forderung $a < c < b$ durch $a \leq c \leq b$ oder durch $a > c > b$ ersetzt.

Wie das Differenzieren (siehe Lemma 35.2) ist auch das Integrieren **linear**.

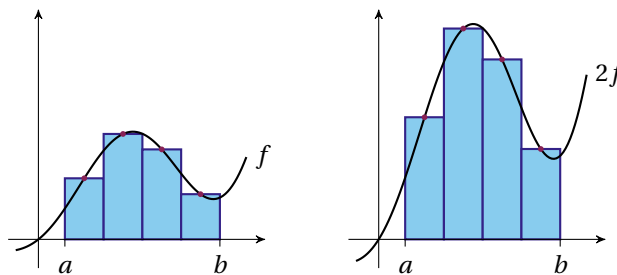
Lemma 36.3

Sind f und g auf $[a, b]$ integrierbar und ist λ eine reelle Zahl, dann sind auch λf und $f + g$ auf $[a, b]$ integrierbar und es gilt

$$\int_a^b \lambda \cdot f(x) dx = \lambda \cdot \int_a^b f(x) dx \text{ und}$$

$$\int_a^b (f(x) + g(x)) dx = \left(\int_a^b f(x) dx \right) + \left(\int_a^b g(x) dx \right).$$

Für λf mit $\lambda = 2$ demonstriert das die folgende Skizze. Für $f + g$ zeichnen Sie sich am besten selbst eine Skizze. (Die beiden Funktionen werden quasi „übereinander gestapelt“.)



Wir benötigen noch einen Hilfsbegriff.

Definition

Ist F eine differenzierbare Funktion, deren Ableitungsfunktion f ist, so nennt man F eine **Stammfunktion** (engl. *antiderivative*) von f .

Z. B. ist $2x^2$ eine Stammfunktion von $4x$. Man beachte, dass der unbestimmte Artikel *eine* verwendet wird. Es gibt nämlich immer unendlich viele Stammfunktionen, weil konstante Terme beim Ableiten verschwinden. $2x^2 - 5$ ist auch eine Stammfunktion von $4x$. Andere Stammfunktionen gibt es aber nicht, wie das folgende Lemma klarstellt.

Lemma 36.4

Ist F eine Stammfunktion von f und $c \in \mathbb{R}$, dann ist $x \mapsto F(x) + c$ ebenfalls eine Stammfunktion von f . Ist G eine andere Stammfunktion von f , dann gibt es eine reelle Zahl d , so dass man G als $G(x) = F(x) + d$ darstellen kann.

Aufgabe 36.6. Geben Sie zwei verschiedene Stammfunktionen von $\sin x$ an.

Aufgabe 36.7. Wir wissen, dass $r x^{r-1}$ die Ableitung von x^r ist. Was ist also die Stammfunktion von x^s ? (Und für welches s funktioniert das nicht?)

Der folgende Satz wird auch **Hauptsatz der Differential- und Integralrechnung** genannt. Der Name deutet bereits die Wichtigkeit an.

Fundamentalsatz der Analysis

- Wenn f eine stetige Funktion auf dem abgeschlossenen Intervall $[a, b]$ ist, dann ist die durch

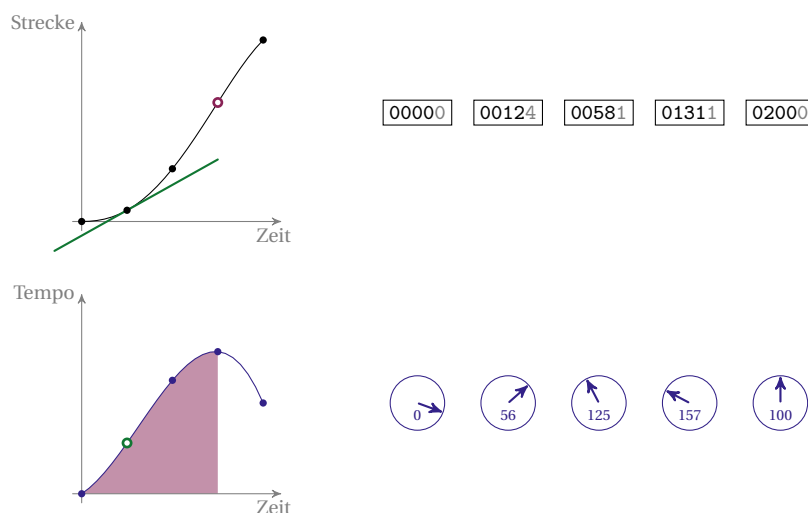
$$F: \begin{cases} [a, b] \rightarrow \mathbb{R} \\ x \mapsto \int_a^x f(t) dt \end{cases} \quad (36.2)$$

definierte Funktion F differenzierbar und eine Stammfunktion von f .

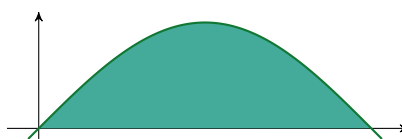
- Wenn f stetig auf $[a, b]$ und G irgendeine Stammfunktion von f ist, dann gilt für alle $x \in [a, b]$ die **Newton-Leibniz-Formel**:

$$\int_a^x f(t) dt = G(x) - G(a)$$

Dieser Satz ist eine der ganz grundlegenden Aussagen der Analysis, weil er einen Zusammenhang zwischen Differenzieren und Integrieren herstellt. Die beiden Operationen sind (unter gewissen Bedingungen) invers zueinander. Man kann sich das z. B. (siehe folgende Skizze) anhand einer Autofahrt klarmachen, bei der man entweder nur den **Tacho** oder nur den **Kilometerzähler** beobachtet. Würde man die entsprechenden Daten für *jeden* Zeitpunkt zur Verfügung haben, so könnte man jeweils eines aus dem anderen rekonstruieren: Die momentane Geschwindigkeit ist die Ableitung der „Streckenkurve“; die aktuelle zurückgelegte Strecke ergibt sich durch Integrieren der „Tempofunktion“.



Insbesondere liefert die Newton-Leibniz-Formel auch einen einfachen Weg, Integrale zu berechnen. (Allerdings unter der Voraussetzung, dass man eine Stammfunktion kennt.) Betrachten wir als Beispiel $\int_0^\pi \sin x \, dx$, siehe Skizze.



numerische Quadratur

Wenn wir diese Fläche mit der „Rechtecksmethode“ annähern (man nennt das übrigens *numerische Quadratur*), erhalten wir je nach Rechengenauigkeit und Schrittweite vielleicht ein Ergebnis wie 2.0000822490709864. (Probieren Sie es selbst mit einem Computer aus!) Für praktische Berechnungen ist so etwas häufig ausreichend, aber man kann sich nie sicher sein, ob das Ergebnis *genau* stimmt. In diesem Fall kennen wir jedoch eine Stammfunktion, nämlich $-\cos x$, und können *symbolisch integrieren*. Nach der **Newton-Leibniz-Formel** ergibt sich

symbolische
Integration

$$\int_0^{\pi} \sin x \, dx = -\cos \pi - (-\cos 0) = -(-1) + 1 = 2. \quad (36.3)$$

Das ist das *exakte* Ergebnis und es erfordert wesentlich weniger Rechenaufwand!

Weil man es bei der Berechnung von Integralen oft mit Ausdrücken von der Form $F(b) - F(a)$ zu tun hat, haben sich dafür diverse abkürzende Schreibweisen eingebürgert. Die beiden häufigsten sehen so aus:

$$F(b) - F(a) = F(x) \Big|_{x=a}^{x=b} = \left[F(x) \right]_{x=a}^{x=b}.$$

Wenn die Variable aus dem Zusammenhang hervorgeht (was meistens der Fall ist), lässt man sie oben und unten auch einfach weg. Gleichung (36.3) könnte dann z. B. so aussehen:

$$\int_0^{\pi} \sin x \, dx = -\cos x \Big|_0^{\pi}.$$

Aufgabe 36.8. Machen Sie sich durch eine Skizze klar, was die Funktion F aus (36.2) berechnet.

Aufgabe 36.9. Überzeugen Sie sich, dass man in (36.3) dasselbe Ergebnis erhält, wenn man eine andere Stammfunktion wie z. B. $x \mapsto 3 - \cos x$ verwendet.

Aufgabe 36.10. Berechnen Sie $\int_1^3 (x^2 + x + 1) \, dx$ und $\int_0^2 2^x \, dx$.

Aufgabe 36.11. Berechnen Sie $\int_1^3 \frac{dx}{x}$ und $\int_{-\pi}^{\pi} \sin x \, dx$. (Die weit verbreitete Schreibweise $\frac{dx}{x}$ ist eine Abkürzung für $\frac{1}{x} dx$.)

Aufgabe 36.12. Berechnen Sie $\int_{\pi/4}^{\pi/2} \frac{2dx}{\tan x}$. (Hinweis: Differenzieren Sie $\ln \sin x$.)

Die entscheidende Einschränkung beim symbolischen Integrieren ist, dass man eine Stammfunktion finden muss. Während das Differenzieren von **elementaren Funktionen** ein simpler mechanischer Prozess ist, der immer zum Erfolg führt, ist Integrieren schwierig und häufig nicht einmal möglich – jedenfalls nicht in dem Sinne, dass eine Stammfunktion sich elementar hinschreiben lässt. Und zwar selbst

dann, wenn nach dem [Fundamentalsatz](#) eine Stammfunktion existieren muss.²⁰ In der Praxis wird daher häufig auf die numerische Quadratur zurückgegriffen.

Bemerkungen

In der Schule wird leider manchmal das Kinderwort „Aufleiten“ für das Integrieren verwendet. Abgesehen davon, dass das ein bisschen dummlich klingt, ist es didaktisch falsch, weil damit der Eindruck erzeugt wird, Integrieren sei als die Umkehrung des Ableitens definiert. Das ist aber falsch und so kann man den [Fundamentalsatz](#) auch nicht interpretieren! Integrieren ist unabhängig vom Differenzieren das Berechnen von Flächen bzw. kontinuierliches Summieren, so wie es am Anfang von Abschnitt 36 eingeführt wird. Der Fundamentalsatz besagt lediglich, dass man *unter gewissen Umständen* Integrale mithilfe von Stammfunktionen berechnen kann.

Bei der Erklärung des Wortes *summa* sollte klar geworden sein, dass der Integrand mit dx multipliziert wird (wobei man den Multiplikationspunkt nie hinschreibt). Daher müssen wie bei der normalen Multiplikation wegen der Regel „Punkt- vor Strichrechnung“ ggf. Klammern gesetzt werden. Für die Integranden $2x$ und $2 + x$ schreibt man z. B.

$$\int_a^b 2x \, dx, \quad \text{aber} \quad \int_a^b (2 + x) \, dx.$$

Alle Stammfunktion einer Funktion sind gleichberechtigt. $2x^2$ ist beispielsweise nicht irgendwie „besser“ als $2x^2 - 5$, weil die erste Funktion einfacher aussieht.

★ Differenzieren und Integrieren im Computer

Eine ganz einfache „selbstgestrickte“ Funktion für die numerische Quadratur könnte folgendermaßen aussehen:

```
def area(f, a, b, n):
    step = (b - a) / n
    sum = 0
    for i in range(n):
        sum += f(a + step/2 + i * step)
    return sum * step
```

Für das Lernen und Verstehen reicht das schon aus. Für die Praxis sollte man eher auf Funktionen wie [quad](#) aus der Bibliothek [SCIPY](#) zurückgreifen.

Die schon mehrfach erwähnte Bibliothek [SYMPY](#) kann symbolisch integrieren (und auch differenzieren). Mehr zu diesen Themen im Buch [Konkrete Mathematik \(nicht nur\) für Informatiker](#). Speziell zum Differenzieren verweise ich noch einmal auf das Video [Wie man Computern das Ableiten beibringt](#), das demonstriert, wie man das Differenzieren mit vergleichsweise geringem Aufwand auch selbst programmieren kann.

²⁰Im Video [Warum man manche Funktionen nicht integrieren kann](#) wird das ausführlich erklärt.

VIII

Stochastik

Stochastik ist in der deutschen Sprache der Oberbegriff für die mathematische Disziplin der Wahrscheinlichkeitsrechnung und deren Anwendungen in der Statistik. Das Video [Was ist eigentlich Wahrscheinlichkeit?](#) bietet eine ausführliche Darstellung der Zusammenhänge und Unterschiede zwischen diesen beiden Gebieten und zeigt unter anderem, dass ihre Ergebnisse durchaus kontrovers interpretiert werden, wohingegen Einigkeit über ihre mathematische Behandlung besteht. Das Ziel des Kapitels ist eine erste Einführung in statistische Methoden, die Sie eventuell im weiteren Verlauf Ihres Studiums verwenden werden.

37. Wahrscheinlichkeit

Grundlegende Methoden der Wahrscheinlichkeitsrechnung werden in der Schule thematisiert. Dieser Abschnitt setzt daher voraus, dass Sie zumindest eine intuitive Vorstellung von Wahrscheinlichkeit und eine gewisse Erfahrung im Umgang damit haben. Erreicht werden soll in erster Linie eine Präzisierung und Verallgemeinerung der Grundbegriffe. Wir fangen dafür mit Beispielen an, von denen jedes etwas komplizierter als das vorherige ist.

Im ersten Beispiel betrachten wir das Werfen zweier Würfel. Dieser Vorgang soll mathematisch *modelliert* werden. Das bedeutet einerseits, dass er mit mathematischen Objekten wie Zahlen, [Mengen](#), [Tupeln](#) und so weiter dargestellt wird, und andererseits, dass er in gewissem Sinne vereinfacht wird, indem Aspekte, die für die Analyse unwichtig sind, weggelassen werden. (Das könnte beispielsweise die Farbe der Würfel sein.) Diesen Vorgang nennt man bekanntlich *Abstraktion*.

Naheliegender ist die Darstellung eines Würfelwurfs als [geordnetes Paar](#) (a, b) , wobei a und b für die Augenzahlen des ersten bzw. zweiten Würfels stehen. So ein Paar würde man als *Ergebnis* bezeichnen. Die Menge aller Ergebnisse wird traditionell meistens Ω genannt. In diesem Fall hätten wir also

Ergebnis

$$\Omega = \{(a, b) : a, b \in [6]\} = [6]^2.$$

Ereignis Diese mengentheoretische Darstellung bietet die Möglichkeit, Teilmengen von Ω als *Ereignisse* zu betrachten. Zum Beispiel ist $\{(a, b) \in \Omega : a + b = 7\}$ das Ereignis, dass die Summe der Augenzahlen 7 ergibt. Durch die aus dem ersten Semester bekannten Zusammenhänge zwischen Mengenoperationen und logische Verknüpfungen wird klar, warum sich Mengen gut für die Modellierung von Ereignissen eignen. Sind nämlich A und B Ereignisse (also Teilmengen von Ω), so erhält man unter anderem die folgenden Zusammenhänge:

- $A \cup B$ ist das Ereignis, dass A *oder* B eingetreten ist.
- $A \cap B$ ist das Ereignis, dass A *und* B eingetreten sind.
- Sind A und B **disjunkt**, so schließen sich A und B gegenseitig aus.
- $A \subseteq B$ bedeutet, dass B eine logische Konsequenz von A ist.
- $\Omega \setminus A$ ist das Gegenteil des Ereignisses A .

Komplement A^c

Da man die mengentheoretische Differenz $\Omega \setminus A$ häufig benötigt, verwendet man dafür die abkürzende Schreibweise A^c und nennt diese Menge das *Komplement* von A (bezüglich Ω). Natürlich ergibt so eine Schreibweise nur Sinn, wenn klar ist, welche Grundmenge (in diesem Fall Ω) gemeint ist.

Aufgabe 37.1. Im obigen Beispiel sei A das Ereignis, dass die Augenzahl des ersten Würfels gerade ist, und B das Ereignis, dass die Augenzahl des zweiten Würfels größer als 3 ist. Stellen Sie A und B in korrekter beschreibender Mengenschreibweise dar und geben Sie die Mächtigkeiten dieser Mengen an.

Aufgabe 37.2. Geben Sie für die Ereignisse aus Aufgabe 37.1 $A \cap B$, $A \cup B$ und A^c sowie deren Mächtigkeiten an. Sind A und B disjunkt? Gilt $A \subseteq B$ oder $B \subseteq A$? Überlegen Sie sich jeweils, wie Ihre Antworten im oben genannten Sinne zu interpretieren sind.

Sinnvollerweise geht man in diesem Würfelbeispiel davon aus, dass alle Ergebnisse dieselbe Wahrscheinlichkeit haben – warum sollte ein Ereignis wie $\{(1, 5)\}$ wahrscheinlicher oder unwahrscheinlicher als $\{(4, 4)\}$ sein? Normiert man noch dadurch, dass man die Wahrscheinlichkeiten 0 und 1 (also 0 bzw. 100 **Prozent**) für Ereignisse festsetzt, die nie oder mit Sicherheit¹ eintreffen, so ergibt sich daraus $P(\{\omega\}) = 1/|\Omega|$ für alle $\omega \in \Omega$, konkret also etwa $P(\{(1, 5)\}) = 1/36$. (Der Buchstabe P , in manchen Büchern auch in der Schriftart $\mathbb{P}$, steht dabei wahlweise für Wörter wie *probability*, *probabilité*, *probabilitas* und so weiter.)

Zudem klingt es vernünftig, dass die Wahrscheinlichkeit eines Ereignisses $A \cup B$ die Summe der Wahrscheinlichkeiten von A und B ist, wenn A und B sich gegenseitig ausschließen (wenn die Mengen also disjunkt sind). Im Beispiel würde daraus unter anderem

$$\begin{aligned} P(\{(5, 6), (6, 5), (6, 6)\}) &= P(\{(5, 6)\} \cup \{(6, 5)\} \cup \{(6, 6)\}) \\ &= P(\{(5, 6)\}) + P(\{(6, 5)\}) + P(\{(6, 6)\}) = 3/36 = 1/12 \end{aligned}$$

folgen und allgemein

$$P(A) = |A|/|\Omega| \tag{37.1}$$

¹Diese Interpretation wird später noch relativiert werden, siehe Fußnote 2 auf Seite 287.

als Wahrscheinlichkeit für beliebige Ereignisse $A \subseteq \Omega$. Insbesondere erhält man $P(\emptyset) = 0$ und $P(\Omega) = 1$.

Eine Zuordnung wie (37.1) nennt man *Laplace-Wahrscheinlichkeit*. Sie setzt voraus, dass alle Ergebnisse dieselbe Wahrscheinlichkeit haben und Ω endlich ist.

Laplace-
Wahrscheinlichkeit

Beachten Sie, dass die Funktion P oben durchgehend den *Ereignissen* und nicht den *Ergebnissen* Wahrscheinlichkeiten zugeordnet hat. Im konkreten Fall scheint diese Unterscheidung zwar nicht so wichtig zu sein, Sie sollten sich jedoch gleich daran gewöhnen, dass formal immer die Teilmengen von Ω und *nicht* die Elemente betrachtet werden.

Aufgabe 37.3. Wie kann man die Ereignisse $\emptyset$ und Ω interpretieren?

Aufgabe 37.4. Berechnen Sie die Wahrscheinlichkeiten für die Ereignisse aus den Aufgaben 37.1 und 37.2.

Beachten Sie, dass in jede Modellbildung Annahmen über die modellierte Realität einfließen, die auch falsch sein können und die damit trotz korrekter mathematischer Vorgehensweise zu falschen Schlussfolgerungen führen können. Hier gehen wir zum Beispiel davon aus, dass die beiden Würfel sich gegenseitig nicht beeinflussen – siehe dazu auch das Ende der Lösung von Aufgabe 37.2. Wenn man etwa dem (durchaus verbreiteten) Irrglauben anhängt, dass das Werfen einer Sechs im ersten Wurf die Chance für eine Sechs im zweiten Wurf verringert, dann darf man den Ereignissen $\{(6, 6)\}$ und $\{(6, 1)\}$ nicht dieselbe Wahrscheinlichkeit zuweisen.

Damit kommen wir zum zweiten Beispiel. Laplace-Wahrscheinlichkeiten sind zwar gerade in Einführungen in die Wahrscheinlichkeitstheorie weit verbreitet, aber es sind durchaus auch andere Zuordnungen denkbar. Hat man es etwa mit einem manipulierten Würfel zu tun, so könnte für diesen die Wahrscheinlichkeit für eine Sechs bei 25 Prozent liegen, während alle anderen Augenzahlen dieselbe Wahrscheinlichkeit haben. Um einen Wurf mit diesem Würfel zu modellieren, bietet sich die Menge $\Omega = [6]$ an und wir würden $P(\{6\}) = 1/4 = 25\%$ setzen. Weil weiterhin $P(A \cup B) = P(A) + P(B)$ für disjunkte Ereignisse A und B gelten soll, ergibt sich mit

$$1 = P(\Omega) = P(\{1\}) + P(\{2\}) + \cdots + P(\{6\})$$

zwangsläufig $P(\{n\}) = (1 - 1/4)/5 = 3/20 = 15\%$ für $n \in [5]$.

Aufgabe 37.5. Geben Sie für dieses Beispiel Mengen A und B für die Ereignisse „gerade Augenzahl“ bzw. „Augenzahl kleiner als 3“ an. Geben Sie außerdem die Wahrscheinlichkeiten $P(A)$ und $P(B)$ an. Wie würden $P(A)$ und $P(B)$ bei einem fairen Würfel aussehen?

Für das nächste Beispiel benötigen wir noch einen Begriff. Dafür sei daran erinnert, dass Addition zunächst nur für zwei Zahlen (und damit wegen der *Assoziativität* für endlich viele) definiert war und wir *einen gewissen Aufwand* betreiben mussten, um zu klären, was mit „unendlichen Summen“ gemeint sein soll. Ebenso können

mit den **Junktoren** $\wedge$ und $\vee$ nur endlich viele Aussagen verknüpft werden, wohingegen die **Quantoren** $\forall$ und $\exists$ in gewissem Sinne **unendliche Konjunktionen bzw. Disjunktionen** sind. Wir vollführen den Schritt hin zum Unendlichen nun auch für die endlichen Mengenoperationen **Vereinigung und Durchschnitt**.

Definition

Ist $\mathcal{A}$ eine Menge von Mengen, dann definieren wir die **Vereinigungsmenge** von $\mathcal{A}$ als

$$\bigcup \mathcal{A} = \bigcup_{M \in \mathcal{A}} M = \{x : \exists M \in \mathcal{A} \, x \in M\}.$$

Ist $\mathcal{A}$ nicht leer, dann definieren wir außerdem die **Schnittmenge** von $\mathcal{A}$ als

$$\bigcap \mathcal{A} = \bigcap_{M \in \mathcal{A}} M = \{x : \forall M \in \mathcal{A} \, x \in M\}.$$

Die Voraussetzung im zweiten Teil, dass $\mathcal{A}$ nicht leer ist, ist wichtig. Würden wir $\bigcap \emptyset$ erlauben, so würden wir so etwas wie „die Menge *aller* Objekte“ erhalten. Das führt auf einen der am Ende von Abschnitt 2 erwähnten Widersprüche.

Wir haben es hier mit drei „Stufen“ zu tun: Objekte, Mengen von Objekten und Mengen von solchen Mengen. Aus didaktischen Gründen verwenden wir in diesem Abschnitt unterschiedliche Symbole für diese Stufen: Kleinbuchstaben, Großbuchstaben und kalligraphische Buchstaben. Diese Unterscheidung werden wir jedoch später wieder aufgeben und sie wird auch nicht in allen Büchern so gehandhabt.

x ist nach dieser Definition ein Element der Vereinigungsmenge $\bigcup \mathcal{A}$, wenn x Element mindestens einer Menge aus $\mathcal{A}$ ist. Ist $\mathcal{A}$ endlich, so ergibt sich eigentlich nichts Neues, z. B.

$$\begin{aligned} \bigcup \{A\} &= \{x : \exists X \in \{A\} \, x \in X\} = \{x : x \in A\} = A, \\ \bigcup \{A, B\} &= \{x : \exists X \in \{A, B\} \, x \in X\} = \{x : x \in A \vee x \in B\} = A \cup B, \\ \bigcup \{A, B, C\} &= \{x : \exists X \in \{A, B, C\} \, x \in X\} \\ &= \{x : x \in A \vee x \in B \vee x \in C\} = A \cup B \cup C \end{aligned}$$

und so weiter. Für die Schnittmenge $\bigcap \mathcal{A}$ sieht das analog ebenso aus.

Interessant wird es erst, wenn $\mathcal{A}$ unendlich ist. Und dafür sollte man sich an den Absatz vor Lemma 3.1 erinnern. In diesem Sinne kann man sich Vereinigungs- und Schnittmenge informell als

$$\begin{aligned} \bigcup \{A_1, A_2, A_3, \dots\} &= A_1 \cup A_2 \cup A_3 \cup \dots \\ \bigcap \{A_1, A_2, A_3, \dots\} &= A_1 \cap A_2 \cap A_3 \cap \dots \end{aligned} \tag{37.2}$$

vorstellen, obwohl dieser Vergleich etwas hinkt, weil erstens die rechten Seiten dieser beiden „Gleichungen“ nicht definiert sind und weil zweitens die in der obigen Definition auftretende Menge $\mathcal{A}$ nicht notwendig in der Form A_1, A_2 und so weiter „sortiert“ sein muss.

Für eine Vereinigungsmenge wie in (37.2) werden wir in Zukunft auch Schreibweisen wie

$$\bigcup_{n=1}^{\infty} A_n$$

$$\bigcup_{n=1}^{\infty} A_n$$

verwenden und analog auch für die Schnittmenge. Ebenso schreiben wir beispielsweise $\bigcup_{k=1}^3 B_k$ für $B_1 \cup B_2 \cup B_3$ und so weiter. Das sollte Sie an die Notation für Summen erinnern und hoffentlich selbsterklärend sein.

Aufgabe 37.6. Sei A_n die Menge der Teiler von n , also $A_n = \{k \in \mathbb{N} : k \mid n\}$. Geben Sie die Menge $\bigcap_{n=2}^4 A_{2n}$ an.

Aufgabe 37.7. Sei A_n wie in Aufgabe 37.6 definiert. Um welche Menge handelt es sich dann bei $\bigcap_{n=1}^{\infty} A_n$?

Aufgabe 37.8. Sei $\mathcal{A} = \{[n] : n \in \mathbb{N}^+\}$. Geben Sie $\bigcup \mathcal{A}$ und $\bigcap \mathcal{A}$ an. Was ändert sich, wenn wir $\mathbb{N}^+$ durch $\mathbb{N}$ ersetzen?

Aufgabe 37.9. Sei $\mathcal{A} = \{\mathbb{N} \setminus \{n\} : n \in \mathbb{N}\}$. Geben Sie $\bigcup \mathcal{A}$ und $\bigcap \mathcal{A}$ an.

Aufgabe 37.10. Was ist $\bigcup \emptyset$?

Aufgabe 37.11. Geben Sie die folgenden Mengen an:

$$A_1 = \bigcap_{n=1}^{\infty} [3 - 1/n, 5 + 1/n]$$

$$A_2 = \bigcap_{n=1}^{\infty} (3 - 1/n, 5 + 1/n)$$

$$A_3 = \bigcap_{n=1}^{\infty} [3 + 1/n, 5 + 1/n]$$

$$A_4 = \bigcup_{n=0}^{\infty} \{k \in \mathbb{N} : k \leq n\}$$

$$A_5 = \bigcap_{n=0}^{\infty} \{k \in \mathbb{N} : k \leq n\}$$

$$A_6 = \bigcup_{n=7}^{\infty} \{k \in \mathbb{N} : k \leq n\}$$

$$A_7 = \bigcup_{n=0}^{\infty} [n, n+1)$$

$$A_8 = \bigcup_{n=0}^{\infty} (n, n+1)$$

$$A_9 = \bigcap_{n=0}^{\infty} [n, n+1)$$

$$A_{10} = \bigcap_{n=0}^{\infty} [n, n+1]$$

$$A_{11} = \bigcap_{n=0}^{\infty} \{k \in \mathbb{N} : k > n\}$$

Aufgabe 37.12. Nach den de-morganschen Regeln gelten die beiden folgenden Aussagen für beliebige Mengen A , B_1 und B_2 :

$$A \setminus (B_1 \cap B_2) = (A \setminus B_1) \cup (A \setminus B_2)$$

$$A \setminus (B_1 \cup B_2) = (A \setminus B_1) \cap (A \setminus B_2)$$

Formulieren Sie das so um, dass es auch gilt, wenn man $B_1 \cap B_2$ bzw. $B_1 \cup B_2$ durch die Vereinigung bzw. den Durchschnitt von mehr als zwei Mengen ersetzt. Verwenden Sie dafür die oben eingeführte Schreibweise.

Wie wir in Aufgabe 37.12 gesehen haben, kann man die Regeln von De Morgan auf mehr als zwei Mengen ausweiten. Und sie gelten auch noch, wenn man es mit unendlich vielen Mengen zu tun hat:

$$A \setminus \bigcap_{k=1}^{\infty} B_k = \bigcup_{k=1}^{\infty} (A \setminus B_k)$$

$$A \setminus \bigcup_{k=1}^{\infty} B_k = \bigcap_{k=1}^{\infty} (A \setminus B_k)$$

Obwohl diese Regeln sehr wichtig sind und häufig gebraucht werden, sind sie in gewissem Sinne banal. Von ihrer Richtigkeit kann man sich einfach dadurch überzeugen, dass man laut vorliest, was man eigentlich aufgeschrieben hat. Zu der Menge links vom Gleichheitszeichen in der ersten Gleichung gehören alle Elemente von A , die nicht zu $\bigcap_{k=1}^{\infty} B_k$ gehören. Letzteres bedeutet, dass sie nicht in *allen* B_k enthalten sind; mit anderen Worten: sie sind in mindestens einem B_k *nicht* enthalten. Nichts anderes steht auf der rechten Seite. Die analoge Argumentation für die zweite Gleichung überlasse ich Ihnen.

Aufgabe 37.13. Begründen Sie, warum die beiden folgenden „Definitionen“ von $\mathbb{Z}$ mithilfe von $\mathbb{N}$ *beide falsch* (!) sind:

$$\mathbb{Z} = \{(n, -n) : n \in \mathbb{N}\}$$

$$\mathbb{Z} = \{\{n, -n\} : n \in \mathbb{N}\}$$

Wie könnte man die zweite Variante durch das Hinzufügen eines einzigen Symbols „retten“?

Aufgabe 37.14. Fällt Ihnen eine Menge $\mathcal{A}$ mit den drei Eigenschaften $\mathbb{R} \in \mathcal{A}$, $\bigcup \mathcal{A} \in \mathcal{A}$ und $\bigcap \mathcal{A} \in \mathcal{A}$ ein? [Tipp: Probieren Sie mit einfachen Beispielen herum.]

Nach diesem Einschub zur Mengenlehre zurück zur Wahrscheinlichkeitsrechnung. Das nun folgende Beispiel zeigt, dass die Grundmenge Ω aller Ergebnisse nicht notwendig endlich sein muss. Stellen wir uns dafür vor, dass wir so lange eine (faire) Münze werfen, bis die Seite *Kopf* oben lag. Da wir vorab nicht wissen, wie viele Würfe benötigt werden, ergibt eine Obergrenze für die Anzahl der Versuche keinen Sinn. Es ist zwar offenbar sehr unwahrscheinlich, dass man beispielsweise hundertmal werfen muss, aber es würde sich nicht richtig anfühlen, so einen Fall als *unmöglich* auszuschließen. Ein denkbare Modell wäre $\Omega = \mathbb{N}^+$, wobei die Zahl n dafür steht, dass man n Würfe benötigt.

Wie groß sollte die Wahrscheinlichkeit dafür sein, dass man genau n Würfe benötigt? Bezeichnen wir die Münzseite *Kopf* mit 1 und die Seite *Zahl* mit 0, so können wir eine Sequenz von n Würfeln als n -Tupel mit Komponenten aus $\{0, 1\}$ darstellen. Nach Lemma 15.2 gibt es 2^n solcher n -Tupel und nur eines von denen besteht aus $n-1$ Nullen gefolgt von einer einzigen Eins am Ende. Da sicher jede Wurfsequenz dieselbe Wahrscheinlichkeit haben sollte, muss $P(\{n\})$ also den Wert $1/2^n$ haben. Nach Satz 32.5 gilt außerdem

$$\sum_{i=1}^n P(\{i\}) = \sum_{i=1}^n \frac{1}{2^i} = 1.$$

Andererseits muss natürlich $P(\Omega) = 1$ gelten und wir haben $\Omega = \bigcup_{n=1}^{\infty} \{n\}$, wobei die Mengen $\{1\}$, $\{2\}$ und so weiter paarweise disjunkt sind. Es scheint also sinnvoll zu sein, dass $\sum_{i=1}^n P(\{n\})$ dasselbe wie $P(\bigcup_{n=1}^{\infty} \{n\})$ ist. Mit anderen Worten ist es oft hilfreich, wenn man die weiter oben formulierte Forderung, dass die Gleichung $P(A \cup B) = P(A) + P(B)$ für disjunkte Ereignisse A und B gelten sollte, auf unendlich viele Ereignisse ausdehnt. Ein weiteres Beispiel dafür zeigt die folgende Aufgabe.

Aufgabe 37.15. Geben Sie im obigen Modell das Ereignis „Es wird eine gerade Anzahl von Würfeln benötigt“ als Menge an und berechnen Sie die Wahrscheinlichkeit für dieses Ereignis. (Hinweis: Zerlegen Sie die Menge in lauter **Singletons** und bilden Sie dann wie eben den Wert einer Reihe.)

Im letzten Beispiel betrachten wir die in der Praxis häufig auftretende Situation, dass die Ergebnisse messbare Größen wie etwa Längen, Temperaturen oder Gewichte sind. Es liegt nahe, als Menge Ω ein geeignetes **Intervall** zu wählen. Dabei kommt es jedoch zu grundsätzlichen Problemen. Will man beispielsweise jeder Zahl im Intervall $[0, 1]$ dieselbe Wahrscheinlichkeit p zuordnen, ausgewählt zu werden, so kann man nur $p = 0$ wählen. Wäre nämlich p positiv, so könnte man $n = \lceil 1/p \rceil + 1$ setzen und – weil $[0, 1]$ unendlich viele Elemente hat – n verschiedene Zahlen x_1 und x_n aus $[0, 1]$ auswählen, für die dann

$$P(\{x_1, \dots, x_n\}) = P(\{x_1\}) \cup \dots \cup P(\{x_n\}) = n \cdot p > 1$$

gelten würde.

Die Setzung $P(\{x\}) = 0$ für alle $x \in [0, 1]$ ist zumindest seltsam, weil ja eine dieser Zahlen gewählt werden muss. Man darf also eine Wahrscheinlichkeit von 0 nicht als *unmöglich* interpretieren.² Während das jedoch „nur“ ungewohnt ist, gibt es ein weiteres Problem, das mathematisch wesentlich ernsthafter ist. In den vorherigen Beispielen konnte man jede Teilmenge von Ω als Ereignis betrachten; bei einem Intervall ist das jedoch nicht zielführend. Man kann beweisen, dass es unmöglich ist, auf einem Intervall wie $[0, 1]$ etwas zu definieren, was sich wie eine Wahrscheinlichkeit verhält, keine Zahl „bevorzugt“ und gleichzeitig allen Teilmengen einen Wert zuordnet.³ Daher betrachtet man in der Wahrscheinlichkeitstheorie im Allgemeinen nicht alle Teilmengen, sondern nur eine Auswahl.

Ist Ω eine nichtleere Menge und $\mathcal{A}$ eine Teilmenge von $\mathcal{P}(\Omega)$ (also eine Menge von Teilmengen von Ω), so bezeichnet man $\mathcal{A}$ als **σ -Algebra** auf Ω , wenn die folgenden drei Bedingungen erfüllt sind:

- (i) $\Omega \in \mathcal{A}$
- (ii) Aus $A \in \mathcal{A}$ folgt $A^c = \Omega \setminus A \in \mathcal{A}$.
- (iii) Sind A_1, A_2 und so weiter Elemente von $\mathcal{A}$, dann ist auch $\bigcup_{n=1}^{\infty} A_n$ ein Element von $\mathcal{A}$.

Definition

²Im Fachjargon spricht man von *fast unmöglich*.

³Schauen Sie dazu bei Interesse [das Video über Vitali-Mengen](#) an.

fast unmöglich

Bevor wir uns überlegen, was aus dieser Definition folgt, schauen wir uns zunächst ein simples Beispiel an. Dafür sei Ω wie oben die Menge der Würfelerggebnisse, also $\Omega = [6]$. $\mathcal{P}(\Omega)$ ist dann eine σ -Algebra. Es ist offensichtlich, dass alle drei Bedingungen erfüllt sind.⁴ Und das hat nichts mit irgendwelchen speziellen Eigenschaften von Ω zu tun. Für jede nichtleere Menge D ist $\mathcal{P}(D)$ eine σ -Algebra auf D . Es gibt noch ein Beispiel, das „immer klappt“: Ist D irgendeine nichtleere Menge, so ist $\{\emptyset, D\}$ eine σ -Algebra auf D .

Aufgabe 37.16. Überzeugen Sie sich durch explizites Überprüfen aller drei Bedingungen, dass sowohl $\mathcal{P}(D)$ als auch $\{\emptyset, D\}$ wirklich σ -Algebren auf D sind.

Wenn das schon alles wäre, wäre die Definition natürlich sinnlos. Daher nun ein Beispiel für eine nichttriviale σ -Algebra auf $\Omega = [6]$. Wir definieren:

$$\mathcal{A} = \{A \subseteq \Omega : |A \cap \{1, 2\}| \in \{0, 2\}\} \quad (37.3)$$

Die Schreibweise $\{A \subseteq \Omega : \dots\}$ ist eine Abkürzung für $\{A \in \mathcal{P}(\Omega) : \dots\}$.

Aufgabe 37.17. Beschreiben Sie mit Ihren eigenen Worten, welche der Teilmengen von Ω zu $\mathcal{A}$ gehören.

Aufgabe 37.18. Überprüfen Sie, dass $\mathcal{A}$ aus (37.3) tatsächlich eine σ -Algebra auf der Menge $[6]$ ist.

Bedingung (ii) sorgt dafür, dass die Definition einer σ -Algebra in gewissem Sinne symmetrisch ist. Aus (i) und (ii) folgt nämlich sofort, dass $\emptyset = \Omega \setminus \Omega$ zu jeder σ -Algebra gehört. Man hätte Bedingung (i) also durch die Bedingung $\emptyset \in \mathcal{A}$ ersetzen können.⁵ Ebenso könnte man in Bedingung (iii) die Vereinigung durch den Durchschnitt ersetzen; das folgt aus (ii) und den Regeln von De Morgan.

Aufgabe 37.19. Ist in den folgenden Fällen jeweils $\mathcal{A}$ eine σ -Algebra auf Ω ? Begründen Sie Ihre Antworten.

- (i) $\Omega = [6]$, $\mathcal{A} = \{A \subseteq \Omega : |A| > 0\}$
- (ii) $\Omega = [6]$, $\mathcal{A} = \{A \subseteq \Omega : |A| \text{ ist gerade}\}$
- (iii) $\Omega = \mathbb{N}$, $\mathcal{A} = \{\emptyset, \mathbb{N}, \{42\}, \mathbb{N} \setminus \{42\}\}$
- (iv) $\Omega = \mathbb{N}$, $\mathcal{A} = \{A \subseteq \Omega : A \text{ ist unendlich}\}$
- (v) $\Omega = \mathbb{N}$, $\mathcal{A} = \{A \subseteq \Omega : A \text{ ist endlich oder } A^c \text{ ist endlich}\}$

Die für die Wahrscheinlichkeitsrechnung wichtigste σ -Algebra ist die sogenannte **borelsche σ -Algebra** $\mathcal{B}(\mathbb{R})$. Dabei handelt es sich um eine σ -Algebra auf $\mathbb{R}$. (Und man kann analog solche σ -Algebren auch für andere Grundmengen wie Intervalle,

⁴Die einzige Bedingung, bei der es evtl. zu Missverständnissen kommen könnte, ist die dritte. Daher sei explizit darauf hingewiesen, dass in (iii) nicht die Rede davon ist, dass die Mengen A_n alle verschieden sein müssen. Es wird also insbesondere auch gefordert, dass die Vereinigung von endlich vielen Elementen von $\mathcal{A}$ wieder ein Element von $\mathcal{A}$ sein muss.

⁵Denn daraus würde umgekehrt zusammen mit (ii) wieder folgen, dass Ω zu $\mathcal{A}$ gehört.

$\mathbb{R}^2$, $\mathbb{R}^3$ oder den Einheitskreis angeben.) Ich möchte nicht im Detail darauf eingehen, wie diese σ -Algebra aussieht,⁶ sondern nur die für uns wesentlichen Fakten aufzählen:

- Alle Mengen von reellen Zahlen, mit denen wir es in der Praxis zu tun haben, gehören zu $\mathcal{B}(\mathbb{R})$.
- Die meisten Teilmengen von $\mathbb{R}$ gehören jedoch *nicht* zu $\mathcal{B}(\mathbb{R})$. (Dabei ist die Formulierung „die meisten“ in einem präzisierbaren mathematischen Sinne gemeint.)
- Trotzdem ist es erstaunlicherweise unmöglich, eine konkrete Menge, die nicht zu $\mathcal{B}(\mathbb{R})$ gehört, explizit anzugeben.
- Die borelsche σ -Algebra löst das in Fußnote 3 auf Seite 287 erwähnte Problem der „nicht messbaren“ Mengen. In $\mathcal{B}(\Omega)$ liegen alle Mengen, die man für die Wahrscheinlichkeitsrechnung braucht, während die „problematischen“ Mengen nicht dazugehören.

Sie finden es vielleicht ein bisschen enttäuschend, dass wir den technischen Aufwand der Definition der σ -Algebren betrieben haben, um dann am Ende lapidar zu sagen, dass es bei den für uns relevanten Mengen schon irgendwie gehen wird. Mir war es jedoch wichtig, auf die prinzipiellen technischen Schwierigkeiten hinzuweisen, ohne Sie an dieser Stelle mit zu viel Mathematik zu überfrachten. σ -Algebren sorgen dafür, dass wir die im Rahmen der Wahrscheinlichkeitsrechnung notwendigen mengentheoretischen Verknüpfungen von Ereignissen durchführen können, ohne uns ständig Gedanken darüber machen zu müssen, ob die Resultate überhaupt Ereignisse sind.

Damit haben wir das Rüstzeug für die wichtigste Definition der Wahrscheinlichkeitstheorie.

Sei Ω eine nichtleere Menge, $\mathcal{A}$ eine σ -Algebra auf Ω und P eine Funktion von $\mathcal{A}$ nach $[0, 1]$. Dann nennt man das Tupel $(\Omega, \mathcal{A}, P)$ einen **Wahrscheinlichkeitsraum**, wenn die folgenden beiden Bedingungen erfüllt sind:

- (i) $P(\Omega) = 1$
- (ii) Ist $(A_n)_{n \in \mathbb{N}}$ eine Folge paarweise disjunkter $\mathcal{A}$ -Mengen, dann gilt

$$P\left(\bigcup_{n \in \mathbb{N}} A_n\right) = \sum_{n \in \mathbb{N}} P(A_n). \quad (37.4)$$

Die Elemente von $\mathcal{A}$ nennt man **Ereignisse** und die Zahl $P(A)$ nennt man die **Wahrscheinlichkeit** von A .

Definition

Das sind die sogenannten *Kolmogorow-Axiome*, mit denen der junge Russe Andrei Kolmogorow 1933 auf einen Schlag diverse Probleme löste, die sich im 19. Jahrhundert im Zuge der Anwendung der Wahrscheinlichkeitstheorie auf physikalische Fragestellungen ergeben hatten. Mit diesen Axiomen wird präzise definiert, wie

Kolmogorow-Axiome

⁶Einen Eindruck davon bekommt man durch [mein Video über das Lebesgue-Integral](#).

sich etwas, dass sich *Wahrscheinlichkeit* nennt, in der Welt der Mathematik zu verhalten hat, ohne dass die Frage beantwortet wird, was Wahrscheinlichkeit im „wirklichen Leben“ ist. Das ist eine für die moderne Mathematik typische Vorgehensweise.

Eine Fülle von einfachen Beispielen für Wahrscheinlichkeitsräume erhält man, wenn Ω eine nichtleere endliche Menge ist, $\mathcal{A}$ die Potenzmenge von Ω und P die in (37.1) definierte Laplace-Wahrscheinlichkeit. Man spricht dann auch von einem *Laplace-Experiment*.

Laplace-Experiment

Aufgabe 37.20. Modellieren Sie das zufällige Ziehen einer Karte aus einem Kartenspiel mit 52 Karten als Laplace-Experiment. Wie groß ist die Wahrscheinlichkeit, ein Ass zu ziehen? (In so einem Kartenspiel gibt es vier Asse.)

★**Aufgabe 37.21.** Modellieren Sie das zufällige Ziehen von *zwei* Karten aus einem Kartenspiel mit 52 Karten als Laplace-Experiment. Dabei sollen die Karten gemeinsam dem Stapel entnommen werden, man zieht also zwangsläufig zwei verschiedene Karten. Wie groß ist die Wahrscheinlichkeit, zwei Asse zu ziehen? [Hier müssen Sie auf Schulwissen aus der *abzählenden Kombinatorik* zurückgreifen.]

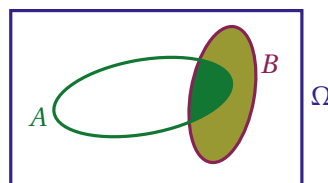
Diverse Eigenschaften von Wahrscheinlichkeiten, die uns ggf. intuitiv ohnehin klar sind, folgen direkt aus den Kolmogorow-Axiomen.

Lemma 37.1

In jedem Wahrscheinlichkeitsraum $(\Omega, \mathcal{A}, P)$ gelten für alle $A, B \in \mathcal{A}$ die folgenden Aussagen:

- (i) $P(\emptyset) = 0$
- (ii) $P(A \cup B) = P(A) + P(B)$, wenn A und B disjunkt sind.
- (iii) $P(A^c) = 1 - P(A)$
- (iv) $P(A) \leq P(B)$, wenn A Teilmenge von B ist.
- (v) $P(A \cup B) = P(A) + P(B) - P(A \cap B)$

Die erste Aussage folgt beispielsweise, indem man in (37.4) $A_n = \emptyset$ für alle $n \in \mathbb{N}$ setzt. Die Reihe auf der rechten Seite kann nach dem *Nullfolgenkriterium* nur konvergieren, wenn die (identischen) Reihenglieder alle verschwinden. Die letzte Aussage kann man sich anhand der folgenden Skizze klarmachen.



Aufgabe 37.22. Das Kartenspiel aus Aufgabe 37.20 hat zwölf „Bilder“ (je vier Könige, Damen und Buben). Außerdem gibt es von jeder „Farbe“ (Kreuz, Pik, Herz, Karo) dreizehn Karten, insbesondere auch einen König, eine Dame und einen Buben von jeder Farbe. Ein Kommilitone von Ihnen sagt: „Die Wahrscheinlichkeit, ein Bild zu ziehen, ist $12/52 = 3/13$. Die Wahrscheinlichkeit, eine Karte der Farbe Herz zu ziehen,

beträgt $1/4$. Also ist die Wahrscheinlichkeit dafür, dass man ein Bild *oder* eine Herz-Karte zieht, $3/13 + 1/4 = 25/52$.“ Wo ist der Denkfehler und was ist die richtige Antwort?

Aufgabe 37.23. Sie würfeln mit drei Würfeln. Wie groß ist die Wahrscheinlichkeit, dass die Gesamtaugenzahl mindestens fünf ist?

Aufgabe 37.24. Weil die Frage an dieser Stelle oft kommt, gebe ich sie an Sie weiter: Ist das in der Lösung zu Aufgabe 37.23 gewählte Modell richtig? Warum muss man z. B. die Würfe $(1, 4, 5)$ und $(1, 5, 4)$ unterscheiden; ist das nicht derselbe Wurf?

Aufgabe 37.25. In einem bekannten Experiment des Psychologen und Nobelpreisträgers **Daniel Kahneman** erhielten die Teilnehmer die folgende Beschreibung einer fiktiven Person: „Linda ist 31 Jahre alt, Single, sehr intelligent und äußert unverblümt ihre Meinung. Sie studierte Philosophie und engagierte sich während ihres Studiums gegen Diskriminierung und soziale Ungerechtigkeiten. Sie nahm auch an Anti-Atom-Demonstrationen teil.“ Sie wurden dann gefragt, welche der beiden folgenden Alternativen sie für wahrscheinlicher hielten:

- (i) Linda ist Bankangestellte.
- (ii) Linda ist Bankangestellte und aktive Feministin.

Die Mehrheit entschied sich für (ii). Warum ist das ein Trugschluss? Verwenden Sie in Ihrer Antwort das in diesem Kapitel entwickelte mathematische Vokabular.

Sie bekommen beim Kartenspielen fünf Karten zugeteilt. Aufgrund irgendeiner illegalen Schummelei wissen Sie, dass zwei Mitspieler, die bereits Karten erhalten haben, beide kein Ass bekommen haben. Dann ist wohl intuitiv klar, dass Ihre Wahrscheinlichkeit, ein oder mehrere Asses zu bekommen, nun höher ist als im allgemeinen Fall. Der folgende Begriff sagt etwas über diese Art von Zusammenhang zwischen verschiedenen Ereignissen aus.

Sei $(\Omega, \mathcal{A}, P)$ ein Wahrscheinlichkeitsraum und seien A und B Ereignisse mit $P(B) \neq 0$. Dann setzt man

$$P(A|B) = \frac{P(A \cap B)}{P(B)} \quad (37.5)$$

und nennt diesen Wert die **bedingte Wahrscheinlichkeit** von A unter der Bedingung B (engl. *conditional probability of A given B*).

Definition

Die (außermathematische) Interpretation ist, dass $P(A|B)$ die Wahrscheinlichkeit für A ist, wenn man weiß, dass B eingetreten ist oder eintreten wird. (Damit ist also nicht notwendig gemeint, dass B zeitlich vor A stattfinden muss. Man sollte sich B eher als Zusatzinformation vorstellen.)

Dass (37.5) sinnvoll ist, wird hoffentlich durch die folgende Grafik klar. Stellen Sie sich vor, Sie werfen mit einem Pfeil auf das Quadrat. A ist das Ereignis, dass Sie den grünen Kreis links treffen, B ist das Ereignis, dass der Pfeil in dem braunen Rechteck rechts landet. Wenn Sie wissen, dass B eintreten wird, dann kann A nur noch eintreten, wenn die blaue Schnittmenge $A \cap B$ der beiden Flächen getroffen

wird. Die Wahrscheinlichkeit dafür ist das Verhältnis des blauen Flächeninhalts zum braunen.



Aufgabe 37.26. Sie werfen zwei Würfel. Wie groß ist die Wahrscheinlichkeit, dass die Augensumme eine **Primzahl** ist, wenn der erste Wurf eine Zwei war?

Aufgabe 37.27. Es werden zwei Würfel geworfen. A sei das Ereignis, dass die Augenzahlen der beiden Würfel verschieden sind, und B das Ereignis, dass die Gesamtanzahl sechs ist. Berechnen Sie $P(A|B)$ und $P(B|A)$.

Aufgabe 37.28. Eine faire Münze wird viermal geworfen. A sei das Ereignis, dass *Zahl* genau dreimal oben liegt. B sei das Ereignis, dass nach *Kopf* nie *Zahl* und vor *Zahl* nie *Kopf* kommt. Ermitteln Sie $P(A)$, $P(B)$, $P(A|B)$ und $P(B|A)$.

Aufgabe 37.29. Beim Münzbeispiel von Seite 286 sei A das Ereignis, dass man mindestens drei Würfe braucht, bis zum ersten Mal *Kopf* geworfen wird. B sei das Ereignis, dass sich der Rest 1 ergibt, wenn man die Anzahl der benötigten Würfe durch 3 dividiert. Berechnen Sie $P(A)$, $P(B)$ und $P(A|B)$.

Wenn sich $P(A)$ und $P(A|B)$ unterscheiden, kann man das so interpretieren, dass das Ereignis B das Ereignis A beeinflusst. Das folgende Konzept formalisiert für beliebig große Mengen von Ereignissen, dass solche Einflüsse *nicht* vorliegen.

Definition

Sei $(\Omega, \mathcal{A}, P)$ ein Wahrscheinlichkeitsraum. Eine Menge $\mathcal{B}$ von Ereignissen heißt **(vollständig) unabhängig**, wenn für jede endliche Teilmenge $\{A_1, A_2, \dots, A_n\}$ von $\mathcal{B}$ die Gleichung

$$P(A_1 \cap A_2 \cap \dots \cap A_n) = P(A_1) \cdot P(A_2) \cdot \dots \cdot P(A_n) \quad (37.6)$$

gilt.

Insbesondere heißt das, dass zwei Ereignisse A und B unabhängig genannt werden, wenn (37.6) mit $n = 2$ gilt, also $P(A \cap B) = P(A)P(B)$. Daraus folgt für den Fall $P(B) \neq 0$ durch Kürzen

$$P(A|B) = \frac{P(A \cap B)}{P(B)} = \frac{P(A)P(B)}{P(B)} = P(A)$$

und ebenso $P(B|A) = P(B)$, wenn $P(A)$ nicht verschwindet. Die Interpretation ist also wie gesagt, dass A und B sich gegenseitig nicht beeinflussen. Dass man etwas über B weiß, erhöht nicht das Wissen über die Wahrscheinlichkeit des Eintretens von A und umgekehrt.⁷

⁷Da man A und B auch vertauschen kann, ist die manchmal verwendete Sprechweise, A sei unabhängig *von* B , übrigens falsch. Die richtige Formulierung ist, dass A und B unabhängig sind.

Aufgabe 37.30. Sind beim **amerikanischen Roulette** die Ereignisse „erstes Dutzend“ (die Zahlen 1 bis 12) und „erste Spalte“ (die Zahlen 1, 4, 7 und so weiter) unabhängig?

Aufgabe 37.31. Es wird mit zwei Würfeln gewürfelt. Berechnen Sie, ob die Ereignisse „Augenzahl des ersten Wurfs ist gerade“ (A) und „Gesamtaugenanzahl ist gerade“ (B) unabhängig sind. Überrascht Sie das Ergebnis?

Aufgabe 37.32. Wir betrachten wieder ein Roulettespiel wie in Aufgabe 37.30, gehen aber der Einfachheit halber davon aus, dass die Bank „fair“ ist: es gibt keine Nullen und die anderen 36 Zahlen haben alle dieselbe Wahrscheinlichkeit $1/36$. Sind die drei Ereignisse „1 to 18“ (A), „red“ (B) und „even“ (C) paarweise unabhängig?

Aufgabe 37.33. In einem verschlossenen Behälter befinden sich vier Karten, die folgendermaßen bedruckt sind:

AAB

ABA

BAA

BBB

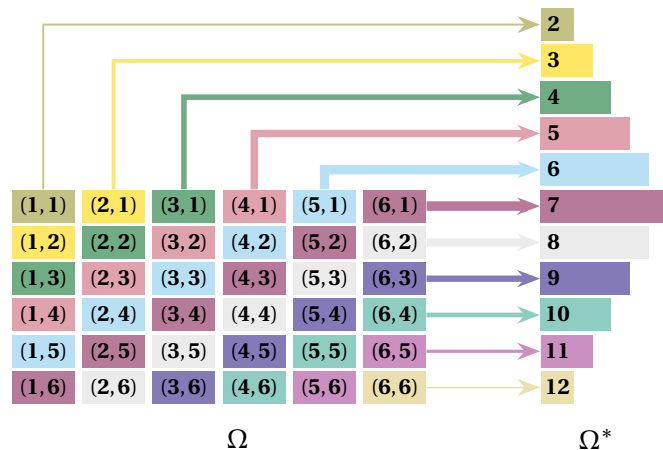
Es wird eine Karte gezogen; die Wahrscheinlichkeiten sollen für alle Karten gleich sein. Für $i = 1, 2, 3$ sei A_i das Ereignis, dass an der i -ten Stelle der gezogenen Karte der Buchstabe A steht. Ist die Menge $\{A_1, A_2, A_3\}$ unabhängig?

38. Zufallsvariablen

In der Praxis arbeitet man typischerweise nicht explizit mit Wahrscheinlichkeitsräumen, die jeweils für ein bestimmtes Problem „maßgeschneidert“ werden, sondern mit „fertig konfektionierten“ sogenannten *Zufallsvariablen*. Dafür schauen wir uns zunächst ein Beispiel an.

Betrachten wir ein Glücksspiel, bei dem zwei Würfel geworfen werden, bei dem es aber nicht um die einzelnen Würfe, sondern lediglich um die Gesamtaugenanzahl geht. Wir *könnten* folgendermaßen vorgehen: Wir überlegen uns, dass die Gesamtaugenanzahl zwischen 2 (zwei Einsen) und 12 (zwei Sechsen) liegen muss. Wir wählen als Ergebnismenge $\Omega^* = \{2, 3, 4, \dots, 12\}$ und berechnen die Wahrscheinlichkeit für jedes Ergebnis, z.B. $P^* (\{11\}) = 1/18$. Wir basteln uns also einen Wahrscheinlichkeitsraum für die spezifische Fragestellung. Für die Berechnung der einzelnen Wahrscheinlichkeiten würden wir dabei wohl das schon bekannte Laplace-Experiment $\Omega = [6]^2$ verwenden. Die Idee der Zufallsvariablen ist nun, gar nicht erst den neuen Raum mit der Ergebnismenge Ω^* zu konstruieren, sondern die Gesamtaugenanzahl als eine *Funktion* der Ereignisse (bzw. der Ergebnisse) eines „Standardraums“ wie Ω zu betrachten.

Solche Funktionen bezeichnet man in der Stochastik üblicherweise mit Großbuchstaben wie X oder Y . In unserem Beispiel wäre X eine Funktion mit Definitionsbereich Ω , die z. B. dem Ergebnis $\omega = (4, 2)$ den Wert $X(\omega) = 6$ zuordnen würde. In der folgenden Skizze zeigen die Pfeile, was X macht. Links unten sehen wir die 36 Elemente von Ω , rechts den Wertebereich von X . Ergebnissen gleicher Farbe wird durch X derselbe Wert zugeordnet.



Der Wertebereich von X ist die Ergebnismenge Ω^* , mit der wir angefangen haben. Wollen wir z.B. die Wahrscheinlichkeit $P^*({11})$ wissen, so können wir den Pfeil, der auf **11** zeigt, zurückverfolgen. Wir sehen, dass die beiden Ergebnisse **(5, 6)** und **(6, 5)** durch X auf **11** abgebildet werden. Die Wahrscheinlichkeit für **{11}** als Ereignis in Ω^* entspricht also der Wahrscheinlichkeit für **{(5, 6), (6, 5)}** als Ereignis in Ω . Entscheidend ist, dass wir diese Wahrscheinlichkeit durch die „Rückprojektion“

$$P^*({11}) = P(\{\omega \in \Omega : X(\omega) = 11\}) = P(\{(5, 6), (6, 5)\})$$

erhalten haben. Bei komplizierteren Wahrscheinlichkeitsräumen kann man sich nicht sicher sein, dass das immer funktioniert, und muss es daher wie in der nun folgenden Definition explizit fordern.

Definition

Ist $(\Omega, \mathcal{A}, P)$ ein Wahrscheinlichkeitsraum und X eine Funktion von Ω nach $\mathbb{R}$, die die Bedingung

$$\{\omega \in \Omega : X(\omega) \leq x\} \in \mathcal{A} \quad (38.1)$$

für alle $x \in \mathbb{R}$ erfüllt, so nennt man X eine (reelle) **Zufallsvariable**. Die Elemente des Wertebereichs von X nennt man die **Realisierungen** der Zufallsvariablen.

Die Axiomatisierung der Wahrscheinlichkeitsrechnung wurde von Kolmogorow ursprünglich in einem deutschsprachigen Lehrbuch veröffentlicht. Dort führte er auch Zufallsvariablen ein, nannte sie aber *Zufallsgrößen*. Obwohl dieser Begriff passender und weniger verwirrend ist, hat sich inzwischen das Wort *Variable* als Rückübersetzung aus dem Englischen etabliert.

Zu beachten ist, dass Zufallsvariablen zunächst einmal einfach *Funktionen* sind. Der Zufall kommt dadurch ins Spiel, dass man mittels (38.1) Wahrscheinlichkeiten dafür angeben kann, dass diese Funktionen bestimmte Funktionswerte annehmen.

Man verwendet für Ereignisse im Raum $(\Omega, \mathcal{A}, P)$, die ähnlich wie in (38.1) entstehen, eine abkürzende Schreibweise, die so funktioniert:⁸

$$\begin{array}{lll} X \leq a & \text{für} & \{\omega \in \Omega : X(\omega) \leq a\} & X \leq a \\ X < a & \text{für} & \{\omega \in \Omega : X(\omega) < a\} & X < a \\ X = a & \text{für} & \{\omega \in \Omega : X(\omega) = a\} & X = a \\ X \in A & \text{für} & \{\omega \in \Omega : X(\omega) \in A\} & X \in A \end{array}$$

Dabei soll a eine reelle Zahl sein und A ein Element von $\mathcal{B}(\mathbb{R})$ (z. B. ein Intervall). Sinngemäß sind auch weitere Schreibweisen wie $a < X \leq b$ zu verstehen. Die Funktion, die jeder Menge $A \in \mathcal{B}(\mathbb{R})$ die Wahrscheinlichkeit $P(X \in A)$ zuordnet, nennt man auch die *Verteilung* von X .

Verteilung

Statt $X \leq a$ findet man in vielen Büchern auch die Schreibweise $\{X \leq a\}$, was leicht missverstanden werden kann. Damit ist natürlich *nicht* die Menge gemeint, deren einziges Element die Formel „ $X \leq a$ “ ist.

Aufgabe 38.1. Sei $(\Omega, \mathcal{A}, P)$ ein Laplace-Experiment, das das dreimalige Werfen einer Münze beschreibt. X soll die Zufallsvariable sein, die angibt, wie oft die Münzseite *Bild* oben lag. Schreiben Sie Ω und X formal korrekt auf und geben Sie die Realisierungen von X an.

Aufgabe 38.2. Wie sieht in der letzten Aufgabe das Ereignis $X < 2$ aus und welchen Wert hat $P(X < 2)$?

Aufgabe 38.3. Informieren Sie sich auf Wikipedia über die Auszahlungsquoten beim *französischen Roulette* und stellen Sie Gewinn bzw. Verlust beim Setzen von fünf Euro auf *Manque* als Zufallsvariable dar. Machen Sie dasselbe auch für *Les trois premiers*.

Aufgabe 38.4. Sei X wie im einführenden Beispiel dieses Abschnitts als Gesamtaugen-zahl zweier Würfelwürfe definiert, allerdings für den gefälschten Würfel von Seite 283. Berechnen Sie $P(X \geq 11)$.

Die entscheidende Aussage über Zufallsvariablen ist die folgende:

Man kann alle⁹ Wahrscheinlichkeiten für eine Zufallsvariable X (also ihre Verteilung) berechnen, wenn man die Werte der Funktion

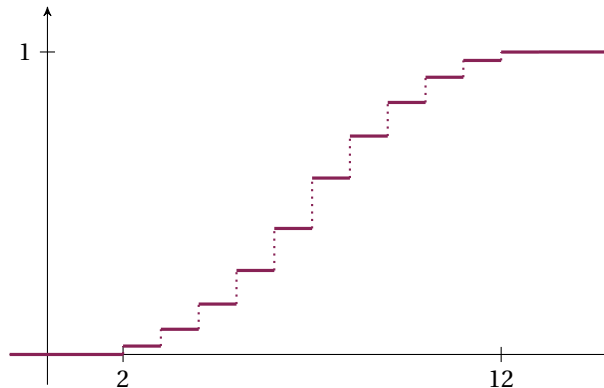
$$F_X: \begin{cases} \mathbb{R} \rightarrow [0, 1] \\ x \mapsto P(X \leq x) \end{cases}$$

kennt, die man die **Verteilungsfunktion** von X nennt. (Im Englischen heißt sie *cumulative distribution function* und wird *CDF* abgekürzt.)

Satz 38.1

⁸Bedingung (38.1) sorgt dafür, dass solche Menge tatsächlich Ereignisse sind. Dabei betrachten wir die Zielmenge $\mathbb{R}$ von X als Wahrscheinlichkeitsraum mit der borelschen σ -Algebra $\mathcal{B}(\mathbb{R})$. Das sind jedoch technische Details, um die wir uns hier nicht kümmern müssen.

Für das Würfelbeispiel vom Anfang des Kapitels sieht F_X folgendermaßen aus:



Wie man an diesem Beispiel sieht, ist F_X im Allgemeinen *nicht stetig*.

Die praktische Bedeutung von Zufallsvariablen ergibt sich mit Satz 38.1 daraus, dass man für das Arbeiten mit X den Definitionsbereich – also den zugrundeliegenden Wahrscheinlichkeitsraum – gar nicht kennen muss, wenn man die Verteilungsfunktion F_X kennt. Ein wesentliches Hilfsmittel dabei ist das folgende Lemma.

Lemma 38.2

Die Verteilungsfunktion F_X einer beliebigen Zufallsvariablen X hat die folgenden Eigenschaften.

- (i) Für $x, y \in \mathbb{R}$ mit $x < y$ gilt $F_X(x) \leq F_X(y)$.
- (ii) Ist (a_n) eine Folge, die bestimmt gegen ∞ divergiert, so konvergiert die Folge $(F_X(a_n))$ gegen 1. Man schreibt das auch als $\lim_{x \rightarrow \infty} F_X(x) = 1$.
- (iii) Ist (a_n) eine Folge, die bestimmt gegen $-\infty$ divergiert, so konvergiert die Folge $(F_X(a_n))$ gegen 0. Man schreibt das auch als $\lim_{x \rightarrow -\infty} F_X(x) = 0$.
- (iv) Ist a eine reelle Zahl und (a_n) eine Folge, die gegen a konvergiert und deren Folgenglieder alle kleiner als a sind, so konvergiert die Folge $(F_X(a_n))$ gegen $P(X < a)$. Man schreibt diesen Grenzwert auch als $\lim_{x \rightarrow a^-} F_X(x)$.
- (v) Für alle $a \in \mathbb{R}$ gilt

$$P(X = a) = P(X \leq a) - P(X < a) = F_X(a) - \lim_{x \rightarrow a^-} F_X(x).$$

Die Wahrscheinlichkeiten aus Punkt (v) entsprechen den Sprungstellen in der obigen Skizze. Wir werden aber *noch sehen*, dass das nicht bei jeder Verteilungsfunktion so aussehen muss.

Wer ein Spielcasino eröffnen will, sollte vorab kalkulieren, wie viel Geld man mit Glücksspielen verdienen kann. Als Beispiel nehmen wir die Zufallsvariable X_M aus Aufgabe 38.3. Da wir davon ausgehen, dass alle Zahlen dieselbe Wahrscheinlichkeit

⁹Mit „alle Wahrscheinlichkeiten“ sind dabei alle von der Form $P(X \in A)$ für $A \in \mathcal{B}(\mathbb{R})$ gemeint. Man kennt also die komplette *Verteilung* von X , wenn man die *Verteilungsfunktion* von X kennt. Das liegt an Bedingung (38.1).

haben, ergibt sich $P(X_M = 5) = 18/37$ und $P(X_M = -5) = 19/37$. Intuitiv würden wir davon ausgehen, dass bei genügend vielen Spielen in ungefähr 18/37 der Fälle der Spieler fünf Euro gewinnt und in 19/37 der Fälle die Bank. Gehen wir der Einfachheit halber von z. B. 3700 Spielen aus, so gewinnt der Spieler davon also circa 1800 und die Bank 1900. In der Bilanz ergibt das $1800 \cdot 5 - 1900 \cdot 5 = -500$, also einen Verlust von 500 Euro für die Spielerseite. Anders ausgedrückt ist davon auszugehen, dass jeder Spieler, der fünf Euro auf *Manque* setzt, *im Durchschnitt* $500/3700$, also $5/37$ Euro bzw. etwa 14 Cent pro Spiel verliert. Den Wert hätte man auch ohne den Umweg über die 3700 Spiele durch

$$-\frac{5}{37} = \frac{18}{37} \cdot 5 + \frac{19}{37} \cdot (-5) = P(X_M = 5) \cdot 5 + P(X_M = -5) \cdot (-5) \quad (38.2)$$

berechnen können. Dieses Konzept wird nun formalisiert.

Ist X eine Zufallsvariable, deren Wertebereich sich als $\{x_i : i \in I\}$ darstellen lässt, wobei $x_i \neq x_j$ für $i \neq j$ gilt und I entweder $\mathbb{N}^+$ oder eine Menge der Form $[n]$ mit $n \in \mathbb{N}^+$ ist, so spricht man von einer **diskreten Zufallsvariable**.

Definition

Solche Wertebereiche (die man übrigens *abzählbare* Mengen nennt) könnten u. a. folgendermaßen aussehen:

abzählbar

- Im Roulettebeispiel gibt es nur die Werte 5 und -5 . Wir hätten $I = [2]$ und $x_1 = 5, x_2 = -5$.
- Handelt es sich bei X um die Gesamtaugenanzahl von zwei Würfeln, so würde man $I = [11]$ und $x_i = i + 1$ für $i \in I$ wählen.
- Zählt X die Anzahl der Würfe, bis bei einer Münze zum ersten Mal *Kopf* oben liegt, dann können wir $I = \mathbb{N}^+$ und $x_i = i$ für alle $i \in I$ nehmen.

Wenn im Folgenden die Ausdrücke I und x_i im Zusammenhang mit diskreten Zufallsvariablen vorkommen, so sind diese wie in der obigen Definition gemeint. Ferner schreiben wir oft abkürzend p_i für $P(X = x_i)$. Die Funktion, die jeder reellen Zahl x die Zahl p_i zuordnet, falls $x = x_i$ gilt, und die anderenfalls den Funktionswert 0 hat, nennt man die *Wahrscheinlichkeitsfunktion* (engl. *probability mass function* bzw. kurz *PMF*) von X .

p_i

Wahrscheinlichkeitsfunktion

Kennt man diese p_i , so weiß man bereits alles über X ,¹⁰ denn die Werte $F_X(x)$ der Verteilungsfunktion erhält man einfach, indem man für die Realisierungen x_i unterhalb von x die zugehörigen Wahrscheinlichkeiten p_i aufsummiert. Im ersten Beispiel des Abschnitts (Summe zweier Würfelwürfe) gilt etwa:

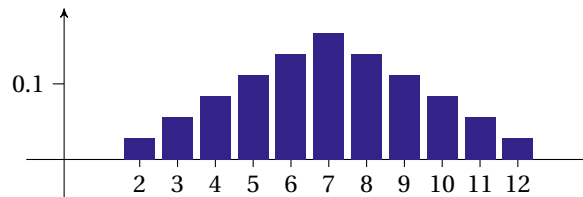
$$F_X(4) = P(X \leq 4) = P(X = 2) + P(X = 3) + P(X = 4)$$

Grafisch kann man sich das so vorstellen, dass man im Funktionsgraphen auf Seite 296 die Sprunghöhen bis inklusive $x = 4$ addiert. Man stellt diskrete Zufallsvariablen aus diesem Grund häufig auch in der Form von *Stab-* oder *Säulendiagrammen* (engl. *bar chart*) dar, bei denen auf der waagerechten Achse die Werte x_i und

Stabdiagramm

¹⁰Das gilt für *diskrete* Zufallsvariablen, *nicht* für alle!

auf der senkrechten Achse die zugehörigen Wahrscheinlichkeiten p_i aufgetragen werden.



Definition

Ist X eine diskrete Zufallsvariable, so bezeichnen wir die Zahl

$$\mathbf{E}(X) = \sum_{i \in I} p_i x_i \quad (38.3)$$

als **Erwartungswert** (engl. *expected value*) von X , wenn sie existiert.¹¹ Man schreibt dafür auch μ_X oder manchmal sogar nur μ .

(38.3) ist offenbar eine Verallgemeinerung von (38.2). Die Idee dahinter ist, dass der Erwartungswert so etwas wie „das durchschnittliche Ergebnis“ ist. (Siehe dazu auch die Lösung von Aufgabe 38.5.)

Nehmen wir als weiteres Beispiel das Werfen der Münze, bis *Kopf* zum ersten Mal erscheint, also $I = \mathbb{N}^+$, $x_i = i$ und $p_i = P(X = x_i) = 1/2^i$ für $i \in I$. Dann ergibt sich (siehe Aufgabe 38.11)

$$\mathbf{E}(X) = \sum_{i \in I} p_i x_i = \sum_{i=1}^{\infty} \frac{i}{2^i} = 2. \quad (38.4)$$

Durchschnittlich muss man also zweimal werfen.

Aufgabe 38.5. Sei X die Augenzahl bei einem Wurf eines fairen Würfels. Geben Sie den Erwartungswert $\mathbf{E}(X)$ an.

Aufgabe 38.6. Machen Sie dasselbe wie in Aufgabe 38.5 für den manipulierten Würfel aus Abschnitt 37.

Aufgabe 38.7. Eine gefälschte Münze wird geworfen, bei der die Wahrscheinlichkeit für *Kopf* 60% beträgt. Ein Spieler setzt fünf Euro und verliert seinen Einsatz, wenn *Kopf* oben liegt. Bei *Zahl* behält er den Einsatz und bekommt fünf Euro dazu. Die Zufallsvariable X soll Gewinn bzw. Verlust des Spielers darstellen. Geben Sie den Erwartungswert von X an.

Aufgabe 38.8. Es wird mit zwei Würfeln gespielt. Ist die Gesamtaugenzahl durch vier teilbar, so erhalten Sie zwanzig Euro. Ergibt sich eine andere Zahl, so müssen Sie den entsprechenden Betrag bezahlen, bei zwei Fünfen also z. B. zehn Euro. Ist das ein faires Spiel, das Sie spielen sollten? Geben Sie dafür für die Zufallsvariable X für Ihren potentiellen Gewinn bzw. Verlust den Erwartungswert an.

¹¹Im Fall $I = \mathbb{N}^+$ handelt es sich um eine **Reihe**, die vielleicht nicht konvergiert. Eine Zufallsvariable muss also nicht notwendig einen Erwartungswert haben.

Aufgabe 38.9. In einem verschlossenen Behälter liegen vier Geldscheine; zwei Fünfeuroscheine, ein Zehneuroschein und ein Zwanzigeuroschein. Es werden zufällig zwei Scheine entnommen. Die Zufallsvariable X stellt den entnommenen Betrag in Euro dar. Geben Sie $E(X)$ an.

Aufgabe 38.10. Mithilfe von Zufallszahlengeneratoren¹² lassen sich viele Ergebnisse der Wahrscheinlichkeitsrechnung experimentell nachvollziehen. Schreiben Sie eine Computersimulation, mit der Sie prüfen können, ob der Wert aus (38.4) realistisch ist.

★**Aufgabe 38.11.** Den Reihenwert in (38.4) habe ich Ihnen einfach vorgesetzt. Man kann ihn jedoch mithilfe des Wissens aus dem [Analysis-Kapitel](#) und etwas Kreativität auch berechnen. Trauen Sie sich das zu?

Verdreifacht man beim Roulette den Einsatz, so verdreifacht sich offensichtlich auch der durchschnittliche Gewinn bzw. Verlust. Und spielt man nebenbei noch ein zweites Glücksspiel, so ergibt sich die zu erwartende durchschnittliche Bilanz des Abends aus der Summe der Erwartungswerte beider Spiele. Zumindest intuitiv sollte das einleuchtend sein. Und es stimmt auch allgemein, weil der Erwartungswert – wie das Ableiten und das Integrieren (siehe Lemma 35.2 bzw. Lemma 36.3) – **linear** ist.

Ist X eine Zufallsvariable und a eine reelle Zahl, so gilt $E(aX) = aE(X)$. Ist Y eine weitere Zufallsvariable mit demselben Definitionsbereich wie X , so gilt $E(X + Y) = E(X) + E(Y)$.

Lemma 38.3

Dabei sind mit aX und $X + Y$ natürlich **Produkte und Summen von Funktionen** gemeint, wie sie in Abschnitt 34 definiert wurden. Die Aussage gilt für beliebige Zufallsvariablen. (Wobei wir hier stillschweigend vorausgesetzt haben, dass die Erwartungswerte existieren. Das werden wir auch im Folgenden machen.) Zumindest für diskrete Zufallsvariablen lässt sie sich leicht nachrechnen, dass Lemma 38.3 korrekt ist; siehe Aufgabe 38.12.

Aufgabe 38.12. Sei $a \in \mathbb{R}$ und X eine Zufallsvariable mit den Realisierungen x_1 bis x_n . Berechnen Sie explizit $E(aX)$ und $aE(X)$, um zu verifizieren, dass die beiden Werte gleich sind. Machen Sie sich klar, dass das genauso funktioniert, wenn X unendlich viele Realisierungen hat.

Aufgabe 38.13. Beim französischen Roulette (Aufgabe 38.3) setzt Frau Müller zehn Euro auf *Les trois premiers* und Herr Müller setzt fünf Euro auf *Manque*. Wenn die Zufallsvariable X den Gesamtgewinn bzw. -verlust von Ehepaar Müller beschreibt, was ist dann der Erwartungswert von X ?

Aufgabe 38.14. Beim Würfeln mit zwei Würfeln sei X die Zufallsvariable, die die Summe der Augenzahlen darstellt, und Y die Anzahl der Würfe mit gerader Augenzahl. Berechnen Sie $E(Y)$, $E(X)E(Y)$ und $E(XY)$, wobei XY natürlich das **Produkt der Funktionen** X und Y sein soll.

¹²Im (zweitiligen) Video [Wie kommt der Zufall in den Computer?](#) lernen Sie evtl. ein paar Dinge, über die Sie bisher nicht nachgedacht haben.

Aufgabe 38.14 hat gezeigt, dass man das Produkt von Zufallsvariablen anders als die Summe nicht einfach mit dem Erwartungswert vertauschen kann. Das klappt im Allgemeinen nur, wenn die Zufallsvariablen *unabhängig* sind.

Definition

Ist $\mathcal{X}$ eine Menge von Zufallsvariablen mit demselben Definitionsbereich, so nennt man diese Menge **unabhängig**, wenn für alle endlichen Teilmengen $\{X_1, \dots, X_n\}$ von $\mathcal{X}$ und alle Teilmengen $\{A_1, \dots, A_n\}$ von $\mathcal{B}(\mathbb{R})$ immer

$$P(X_1 \in A_1)P(X_2 \in A_2) \cdots P(X_n \in A_n) = P\left(\bigcap_{i=1}^n (X_i \in A_i)\right) \quad (38.5)$$

gilt.

Betrachten wir als Beispiel das Werfen zweier Würfel. X_1 sei die Augenzahl des ersten Würfels und X die Summe der beiden Augenzahlen. X_1 und X sind *nicht* unabhängig, weil z. B. die Ereignisse $X_1 = 1$ und $X < 4$ nicht unabhängig sind. (Siehe Aufgabe 38.15.) Ein „klassisches“ Beispiel für zwei unabhängige Zufallsvariablen erhält man hingegen, wenn X_2 die Augenzahl des zweiten Würfels ist und man X_1 und X_2 betrachtet. Sind A und B irgendwelche Teilmengen von $[6]$, so ist $X_1 \in A$ das Ereignis $A \times [6]$ und $X_2 \in B$ ist $[6] \times B$, womit sich als Durchschnitt der beiden das Ereignis $A \times B$ ergibt. Also gilt:

$$\begin{aligned} P(X_1 \in A)P(X_2 \in B) &= P(A \times [6])P([6] \times B) = \frac{|A \times [6]|}{|[6]^2|} \cdot \frac{|[6] \times B|}{|[6]^2|} \\ &= (|A|/6) \cdot (|B|/6) = (|A| \cdot |B|)/36 \\ &= \frac{|A \times B|}{|[6]^2|} = P(A \times B) = P(X_1 \in A, X_2 \in B) \end{aligned}$$

Dabei wurde am Ende die weitverbreitete Konvention verwendet, statt $\cap$ ein Komma zu schreiben.

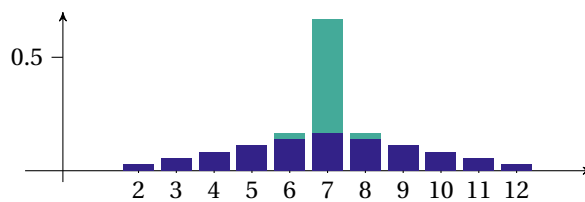
Aufgabe 38.15. Stellen Sie die Ereignisse $X_1 = 1$ und $X < 4$ aus dem letzten Beispiel in der Form $X_1 \in A$ bzw. $X \in B$ dar und verifizieren Sie, dass diese Ereignisse tatsächlich nicht unabhängig sind.

Lemma 38.4

Sind X und Y Zufallsvariablen mit demselben Definitionsbereich, die unabhängig sind, so gilt $E(XY) = E(X)E(Y)$.

Aufgabe 38.16. Beim Wurf zweier Würfel sei Y das Produkt der Augenzahlen. Was ist der Erwartungswert von Y ? Berechnen Sie ihn einmal direkt und einmal mithilfe von Lemma 38.4.

Wir kommen erneut auf das Beispiel mit den zwei Würfeln zurück. Wie bisher sei X die Summe der Augenzahlen. Als Alternative betrachten wir ein Spiel, bei dem vier Seiten eines fairen Würfels mit der Augenzahl 7 beschriftet sind, während auf den restlichen beiden Seiten 6 bzw. 8 steht. Dieser Würfel wird einmal geworfen und seine Augenzahl durch die Zufallsvariable Y beschrieben.



Es ist offensichtlich, dass $E(X) = E(Y) = 7$ gilt. X und Y sind aber ansonsten zwei sehr unterschiedliche Zufallsvariablen. Eine einzige Zahl wie der Erwartungswert liefert uns also nicht alle Informationen über X ; das war ja auch nicht zu erwarten. (Kleines Wortspiel ...) Es gibt jedoch noch weitere Kennzahlen von Zufallsvariablen, die in der Stochastik eine Rolle spielen, und zumindest eine von denen wollen wir uns noch anschauen.

Aufgabe 38.17. Berechnen Sie im obigen Beispiel die Erwartungswerte für $X - E(X)$ und $Y - E(Y)$, wobei mit $X - E(X)$ die Zufallsvariable gemeint ist, die jedem ω den Wert $X(\omega) - E(X)$ zuordnet.

X und Y haben zwar denselben Erwartungswert, man sieht aber an der obigen Grafik deutlich, dass die Werte von X weiter „streuen“. Es bietet sich an, die „durchschnittliche“ *Abweichung* vom Erwartungswert zu untersuchen. Aufgabe 38.17 zeigt jedoch, dass die Zufallsvariable $X - \mu_X$ dafür nicht geeignet ist, weil sich positive und negative Abweichungen gegenseitig neutralisieren. Man könnte stattdessen den Betrag $|X - \mu_X|$ betrachten; allerdings lässt sich mit der Betragsfunktion häufig schlecht arbeiten. Der mathematisch sinnvollste Wert ist der, den wir nun definieren.

Ist X eine Zufallsvariable mit existierendem Erwartungswert μ_X , so nennt man den Wert

$$\text{Var}(X) = E((X - \mu_X)^2)$$

die **Varianz** von X (wenn er existiert). Man schreibt dafür auch σ_X^2 und manchmal nur σ^2 . Die Quadratwurzel der Varianz wird **Standardabweichung** (engl. *standard deviation*) genannt und mit σ_X bezeichnet.

Definition

Da die Varianz von X zwangsläufig eine nichtnegative Zahl sein muss, ist die Standardabweichung immer definiert, wenn die Varianz existiert. Während die Varianz mathematisch leichter zu handhaben ist, hat die Standardabweichung den Vorteil, in derselben Einheit wie die Zufallsvariable angegeben zu werden. Misst X z. B. den Gewinn eines Glücksspiels in Euro, so ist die Varianz von X ein abstrakter Wert in „Quadraturo“; die Standardabweichung ist aber ebenfalls ein Wert in Euro.

Exemplarisch berechnen wir nun Varianz und Standardabweichung der obigen Variable X :

i	1	2	3	4	5	6	7	8	9	10	11
x_i	2	3	4	5	6	7	8	9	10	11	12
p_i	1/36	1/18	1/12	1/9	5/36	1/6	5/36	1/9	1/12	1/18	1/36
$(x_i - \mu_X)^2$	25	16	9	4	1	0	1	4	9	16	25
$p_i(x_i - \mu_X)^2$	25/36	8/9	3/4	4/9	5/36	0	5/36	4/9	3/4	8/9	25/36

Als Summe über die letzte Zeile erhalten wir $\text{Var}(X) = 35/6 \approx 5.8$. Die zugehörige Standardabweichung ist $\sigma_X = \sqrt{35/6} \approx 2.4$. Für Y ergeben sich hingegen die Werte $\text{Var}(Y) = 1/3 \approx 0.3$ und $\sigma_Y = \sqrt{1/3} \approx 0.6$; ein deutliches Zeichen dafür, dass die Werte von X „im Durchschnitt“ weiter vom Erwartungswert entfernt sind als die von Y .

Aufgabe 38.18. Berechnen Sie Varianz und Standardabweichung für die Zufallsvariable aus Aufgabe 38.9.

Aufgabe 38.19. Beim Würfeln mit zwei Würfeln sei X_m die kleinere und X_M die größere der beiden Augenzahlen. (Das ist so gemeint, dass bei einem Wurf wie $\omega = (4, 4)$ die Werte $X_m(\omega) = X_M(\omega) = 4$ angenommen werden.) Berechnen Sie jeweils Erwartungswert und Varianz für diese beiden Zufallsvariablen. Welcher bereits untersuchten Zufallsvariablen entspricht die Summe $X_m + X_M$?

Es folgen ein paar „Rechenregeln“ für die Varianz.

Lemma 38.5

Sei X eine Zufallsvariable, Y eine weitere Zufallsvariable mit demselben Definitionsbereich und $a \in \mathbb{R}$. Dann gelten die folgenden Aussagen:

- (i) $\text{Var}(aX) = a^2 \text{Var}(X)$
- (ii) $\text{Var}(X + a) = \text{Var}(X)$
- (iii) $\text{Var}(X) = \mathbf{E}(X^2) - \mathbf{E}(X)^2$.
- (iv) $\text{Var}(X + Y) = \text{Var}(X) + \text{Var}(Y)$, wenn X und Y unabhängig sind

Ein paar Anmerkungen dazu:

- (i) impliziert, dass die Standardabweichung mit $|a|$ multipliziert wird.
- Dass (ii) stimmt, sollte zumindest anschaulich klar sein, weil das Stabdiagramm von $X + a$ ja nur das um a verschobene Stabdiagramm von X ist.
- (iii) wird manchmal *Verschiebungssatz* genannt.
- Aufgabe 38.19 hat demonstriert, dass (iv) ohne die Forderung der Unabhängigkeit falsch ist.

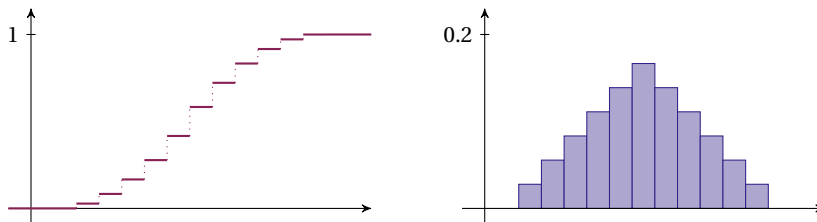
Verschiebungssatz

Aufgabe 38.20. Beim französischen Roulette (siehe Aufgabe 38.3) setzen Sie dreimal nacheinander einen Euro auf *Rouge*. X sei Ihre Bilanz (Gewinn bzw. Verlust) nach diesen drei Spielen. Berechnen Sie die Standardabweichung von X .

★**Aufgabe 38.21.** Alle vier Punkte von Lemma 38.5 folgen mehr oder weniger direkt aus den Rechenregeln für den Erwartungswert. Probieren Sie ruhig einmal aus, das Lemma selbst zu beweisen.

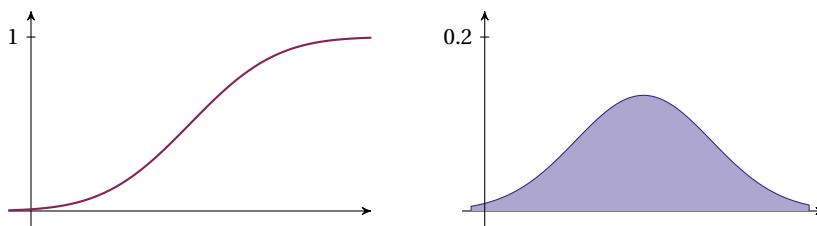
Neben den **diskreten Zufallsvariablen** gibt es eine weitere, für die Praxis der Statistik noch wichtigere Klasse von Zufallsvariablen. Deren Funktionsweise sollen

die beiden folgenden Skizzen klarmachen. In der ersten sieht man links die schon auf Seite 296 abgebildete **Verteilungsfunktion** für die Gesamtaugenzahl zweier Würfel mit den charakteristischen Unstetigkeitsstellen (Sprüngen) bei den Realisierungen der Zufallsvariablen. Rechts daneben ist das Stabdiagramm der einzelnen Wahrscheinlichkeiten p_i zu sehen, die u. a. für die Definition des **Erwartungswertes** verwendet wurden.¹³ (Man beachte die unterschiedlichen Maßstäbe für die vertikale Achse.)



Dabei benötigt man rein technisch nur eins von beiden, weil man die rechte Kurve mittels der linken rekonstruieren kann und umgekehrt: Die Werte rechts entsprechen den Sprunghöhen links.

Stellt man sich nun vor, dass man wesentlich mehr Realisierungen als diese elf hat und dass deren Abstände immer geringer werden, so kommt man im Grenzwert auf die nun folgende Darstellung mit **stetigen** Funktionen.¹⁴



Diese Situation sollte Sie an den **Fundamentalsatz der Analysis** erinnern. Man kann nun die linke Kurve aus der rechten durch **Integrieren** erhalten und die rechte aus der linken durch **Differenzieren**. (Die Grafik darüber könnte man wie bei der Einführung des Riemann-Integrals als grobe Approximation durch Rechtecke interpretieren.)

Auf jeden Fall stellt die Kurve links ebenfalls eine Verteilungsfunktion dar, die alle in Lemma 38.2 aufgeführten Eigenschaften hat. Weil sie aber stetig ist, gibt es keine Sprünge. Nennen wir die zugehörige Verteilung X und wählen wir uns einen beliebigen Punkt $x_0 \in \mathbb{R}$ aus, so gilt nach (v) von Lemma 38.2

$$P(X = x_0) = P(X \leq x_0) - P(X < x_0) = F_X(x_0) - \lim_{x \rightarrow x_0^-} F_X(x)$$

Wegen der Stetigkeit von F_X muss die Differenz rechts verschwinden, weil $F_X(x_0)$ und der Grenzwert identisch sind. Für Zufallsvariablen mit stetiger Verteilungs-

¹³Das könnte man auch als Funktionsgraphen der Wahrscheinlichkeitsfunktion interpretieren, wenn die Stäbe keine Breite hätten.

¹⁴Man sieht hier nur Ausschnitte von Graphen, die eigentlich von $-\infty$ bis ∞ verlaufen.

funktion muss also *jede* Realisierung fast unmöglich sein.¹⁵ Es ergibt nur Sinn, über „Intervall-Wahrscheinlichkeiten“ wie z. B. $P(3 \leq X < 4)$ zu reden.

Definition

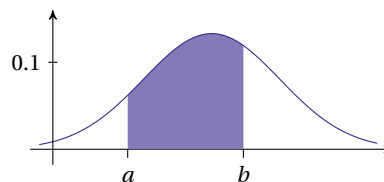
Ein Zufallsvariable X heißt **stetig**, wenn es eine Funktion $f: \mathbb{R} \rightarrow [0, \infty)$ gibt, so dass

$$P(X \leq x) = F_X(x) = \int_{-\infty}^x f(t) dt \quad (38.6)$$

für alle $x \in \mathbb{R}$ und zusätzlich $\int_{-\infty}^{\infty} f(t) dt = 1$ gilt. (Insbesondere müssen all diese Integrale natürlich existieren.) f nennt man dann die **Dichtefunktion** oder einfach die **Dichte** von X (engl. *probability density function* bzw. kurz *PDF*).

Dazu zunächst zwei Anmerkungen. Erstens haben wir offiziell noch gar nicht geklärt, wie die in der Definition vorkommenden sogenannten *uneigentlichen Integrale* zu verstehen sind. Zumindest anschaulich sollte aber klar sein, dass man in (38.6) den Grenzwert von $\int_a^x f(t) dt$ betrachtet, wenn a gegen $-\infty$ geht; man dehnt gewissermaßen die Fläche unter der Kurve nach links bis ins Unendliche aus. Zweitens sei erwähnt, dass es Zufallsvariablen gibt, die weder diskret noch stetig sind. Für die Praxis der Statistik spielen die jedoch keine Rolle.

Entscheidend ist, dass uns die Dichte eine grafische Vorstellung der Wahrscheinlichkeiten liefert. Stellt die Kurve in der folgenden Grafik die Dichte einer Zufallsvariablen X dar, so ist die blau hervorgehobene *Fläche* die Wahrscheinlichkeit $P(a < X < b)$. (Und wegen der Vorüberlegungen spielt es keine Rolle, ob dort $<$ oder $\leq$ steht.)



Die Dichte erlaubt uns außerdem, den Begriff des **Erwartungswertes** sinnvoll zu übertragen.

Definition

Ist X eine stetige Zufallsvariable mit der Dichte f , so bezeichnen wir die Zahl

$$\mathbf{E}(X) = \int_{-\infty}^{\infty} t f(t) dt \quad (38.7)$$

als **Erwartungswert** von X , wenn sie existiert.

Varianz und Standardabweichung werden dann wie im letzten Abschnitt mithilfe des Erwartungswertes definiert und eigentlich ändert sich nichts:

¹⁵Siehe Fußnote 2 auf Seite 287.

Die Lemmata 38.3, 38.4 und 38.5 behalten auch für stetige Zufallsvariablen ihre Gültigkeit.

Lemma 38.6

Die wichtigste Verteilung der Stochastik ist die folgende, die manchmal auch nach **Carl Friedrich Gauß** benannt wird und früher auf dem **Zehn-Mark-Schein** abgebildet war. Ihre Dichtefunktion bildet die typische „Glockenkurve“, wie sie beispielsweise auf Seite 303 zu sehen ist.

Für jede reelle Zahl μ und jede positive Zahl σ wird durch

$$f(t) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left(-\frac{1}{2}\left(\frac{t-\mu}{\sigma}\right)^2\right)$$

eine Dichtefunktion definiert. Zufallsvariablen mit dieser Dichte werden **normalverteilt** genannt und haben den Erwartungswert μ und die Standardabweichung σ . Im Spezialfall $\mu = 0$ und $\sigma = 1$ spricht man von der **Standardnormalverteilung**. Für die **Verteilungsfunktion** der Standardnormalverteilung wird oft der Buchstabe Φ verwendet.

Definition

Um zu verstehen, warum gerade diese Verteilung so wichtig ist, benötigen wir noch ein Konzept, das in der Statistik regelmäßig verwendet wird.

Eine Menge $\mathcal{X}$ von Zufallsvariablen mit demselben Definitionsbereich wird **i. i. d.** genannt, wenn sie **unabhängig** ist und alle Elemente von $\mathcal{X}$ dieselbe Verteilung haben.

Definition

Die Abkürzung i. i. d. steht für *independent and identically distributed*. Ein typisches Beispiel bildet die Menge $\mathcal{X} = \{X_1, X_2\}$ hinter **der Definition der Unabhängigkeit von Zufallsvariablen**. X_1 und X_2 sind *verschiedene* Zufallsvariablen – z. B. gilt $X_1(\omega) = 2$ und $X_2(\omega) = 5$ für $\omega = (2, 5)$ –, aber sie haben *dieselbe* Verteilung, beispielsweise $P(X_1 < 3) = P(X_2 < 3)$ und so weiter.

Das demonstriert nebenbei auch, dass das eigentlich entscheidende Konzept die Verteilungen sind, während die Zufallsvariablen gewissermaßen nur eine „Nebenrolle“ als Träger der Verteilungen spielen. Für bekannte Verteilungen gibt es oft allgemein gebräuchliche Kurzbezeichnungen und man verwendet das Zeichen $\sim$ für „hat die Verteilung“. Z. B. könnte $X \sim \mathcal{N}_{\mu, \sigma^2}$ bedeuten, dass die Zufallsvariable X normalverteilt mit Erwartungswert μ und Standardabweichung σ ist.

Die Bedeutung der Normalverteilung macht nun der folgende Satz klar.

Zentraler Grenzwertsatz von Lindeberg-Lévy

Satz 38.7

Sei $(X_n)_{n \in \mathbb{N}^+}$ eine Folge von i. i. d. Zufallsvariablen mit Erwartungswert μ und Standardabweichung $\sigma > 0$,

$$\overline{X}_n = \frac{1}{n} \sum_{i=1}^n X_i \quad \text{sowie} \quad Z_n = \frac{\overline{X}_n - \mu}{\sigma / \sqrt{n}} \quad (38.8)$$

für $n \in \mathbb{N}^+$. Dann konvergieren die Verteilungsfunktionen der Zufallsvariablen Z_n für $n \rightarrow \infty$ gegen die Verteilungsfunktion Φ der Standardnormalverteilung.

Während in (38.8) die Definition von $\overline{X}_n$ als Mittelwert der ersten n Folgenglieder noch recht einsichtig ist, mag die von Z_n zunächst etwas verwirren. Es geht dabei aber lediglich darum, den Wert $\overline{X}_n$ zu *standardisieren*. Nach Lemma 38.3 hat $\overline{X}_n$ den Erwartungswert μ , also hat der Zähler von Z_n (und damit auch Z_n) den Erwartungswert 0. Und nach Lemma 38.5 hat $\sum_{i=1}^n X_i$ die Varianz $n\sigma^2$ (weil die X_i unabhängig sind) und damit $\overline{X}_n$ die Varianz σ^2/n . Der Zähler von Z_n hat daher ebenfalls die Varianz σ^2/n und der Nenner sorgt dafür, dass die Varianz von Z_n den Wert 1 hat. Der Mittelwert wird also jeweils so verschoben und skaliert, dass Erwartungswert und Varianz gleich bleiben.

Dass dabei im Grenzwert *immer* dieselbe Verteilung herauskommt, ist schon sehr erstaunlich. (Und die Aussage gilt sogar noch [unter schwächeren Voraussetzungen](#).) In Anwendungen kann man das u. a. häufig so interpretieren, dass viele kleine zufällig auftretende Störungen in der Summe eine normalverteilte zufällige Störung ergeben. Das kann etwa der Fehler beim Messen einer physikalischen Größe sein und in diesem Zusammenhang wurde die Normalverteilung auch von Gauß verwendet. Siehe dazu [das Video über die Methode der kleinsten Quadrate](#).

Aufgabe 38.22. Wenn man die Körpergrößen von genügend vielen Personen in Mitteleuropa misst und die Ergebnisse als Stabdiagramm visualisiert, ergibt sich dann näherungsweise eine Glockenkurve?

Die Dichte der Normalverteilung ist übrigens das klassische Beispiel für eine der am Ende von Abschnitt 36 erwähnten Funktionen, die sich nicht elementar integrieren lassen. Die Funktionswerte von Φ sind also grundsätzlich immer Näherungswerte.

Wir haben uns hier nur eine für die Praxis relevante Verteilung näher angeschaut. Weitere wichtige Verteilungen (diskrete und stetige) werden im Buch [Konkrete Mathematik \(nicht nur\) für Informatiker](#) vorgestellt. Dort findet man auch eine wesentlich ausführlichere Einführung in die Stochastik nebst vielen Beispielen dafür, wie man PYTHON für Wahrscheinlichkeitsrechnung und Statistik einsetzen kann.

39. Hypothesentests

schließende Statistik

Damit kommen wir zur sogenannten *schließenden* oder *mathematischen Statistik*. Das Verhältnis derselben zu den Themen der vorherigen Abschnitte lässt sich kurz

folgendermaßen beschreiben: In der Wahrscheinlichkeitsrechnung gehen wir von einer bestimmten Verteilung aus und die entsprechenden Theoreme sagen uns dann, mit welchen Wahrscheinlichkeiten bestimmte Ereignisse eintreten werden, die wir als zu beobachtende *Daten* interpretieren können. In der Statistik sammeln wir hingegen Daten und versuchen, auf dieser Basis auf die Verteilung zu schließen.



Es ist natürlich ganz wichtig, dass es auch in der Statistik um *zufällige* Ereignisse geht; sonst bräuchte man dafür die Resultate der Wahrscheinlichkeitstheorie nicht. Mit statistischen Methoden erzielte Ergebnisse sind also grundsätzlich immer nur mit einer gewissen Wahrscheinlichkeit richtig!

Die wohl am häufigsten verwendete Methode ist die des *Hypothesentests*. Dabei geht es darum, zu einer begründeten Entscheidung zwischen zwei Erklärungen für ein beobachtetes Phänomen zu kommen. Grundsätzlich durchläuft man immer die folgenden Schritte:

Hypothesentest

- (H1) Erstellung eines (wahrscheinlichkeitstheoretischen) Modells
- (H2) Formulierung der Hypothesen und Auswahl des passenden Testverfahrens
- (H3) Festlegung des *Signifikanzniveaus*
- (H4) Erfassung der Daten
- (H5) Ermittlung des Ergebnisses auf der Basis der Daten

Ganz wesentlich für ein korrektes Vorgehen ist, dass diese Schritte in der angegebenen Reihenfolge durchgeführt werden. Hat man nämlich die Daten aus Schritt (H4) bereits und führt erst danach die Schritte (H1) und (H2) durch, so kann man sich nachträglich die Sache so „zurechtbiegen“, dass (H5) die gewünschte Antwort liefert.

Wir gehen das anhand eines konkreten Beispiels durch. Sie vermuten, dass man in Matheklausuren besser abschneidet, wenn man am Abend vorher genügend Hagebuttentee trinkt, und wollen die Korrektheit Ihrer Vermutung durch eine Studie untermauern.

In Schritt (H1) erstellen Sie ein Modell: Aufgrund der Klausuren aus den letzten Jahren gehen Sie davon aus, dass typische Klausurergebnisse *normalverteilt* mit Erwartungswert 10.4 (bezogen auf das Punktesystem der gymnasialen Oberstufe) und Standardabweichung 2.7 sind.

Alle weiteren Schritte basieren darauf, dass das Modell die reale Situation angemessen repräsentiert. An diesem Punkt kann Ihnen die Mathematik nur begrenzt helfen. Im konkreten Beispiel ist die Annahme, dass die Klausurergebnisse normalverteilt sind, durchaus zu vertreten.¹⁶ Es ist jedoch fraglich, ob genügend Datenmaterial vorliegt, um die exakten Werte $\mu = 10.4$ und $\sigma = 2.7$ zu rechtfertigen.

¹⁶Obwohl sie rein formal falsch sein muss, denn dann wären auch „Noten“, die größer als 15 oder negativ sind, möglich (wenn auch ziemlich unwahrscheinlich).

Wir bleiben für das Beispiel bei dieser Annahme, werden aber später noch über Alternativen sprechen, die man einsetzen kann, wenn solche Werte nicht vorliegen.

Außerdem gehen die meisten Testverfahren davon aus, dass die einzelnen Ereignisse (in diesem Fall die Klausurnoten) **unabhängig** sind. Diese Voraussetzung wird damit auch Teil unseres Modells. Aber ist das wirklich so? (Wenn beispielsweise zwei Studierende zusammen gelernt haben, machen sie vielleicht dieselben Fehler.) Wir müssen hoffen, dass die Annahme der Unabhängigkeit eine Abstraktion ist, die nicht zu stark von der Realität abweicht.

In Schritt (H2) werden nun zwei sich widersprechende Hypothesen aufgestellt, die spezielle Namen haben. In unserem Beispiel ist

- | | |
|---------------------|--|
| Nullhypothese | • die <i>Nullhypothese</i> , dass Hagebuttentee am Ergebnis nichts ändert oder es schlimmstenfalls sogar verschlechtert, |
| Alternativhypothese | • während die <i>Alternativhypothese</i> besagt, dass die Zensuren durch den Tee besser werden. |

Das, was Sie eigentlich belegen wollen, ist immer die Alternative. Warum das so gemacht wird und nicht andersherum, klären wir noch.

Ihr Plan ist, eine Klausur mit 20 Studierenden zu schreiben, die alle am Vorabend mindestens einen Liter Hagebuttentee zu sich genommen haben. Sie hoffen natürlich, dass die Ergebnisse besser als üblich sind, um so Ihre These zu stützen. Die Frage ist: Wie viel besser müssen sie sein, um eine verlässliche Aussage machen zu können? Und was genau wäre dann die Aussage?

Sie informieren sich darüber, welches Testverfahren in dieser Situation sinnvoll ist, und finden den sogenannten **Gauß-Test**. Dessen Arbeitsweise soll nun erklärt werden. (In der Praxis werden Sie die mathematische Herleitung überspringen, siehe weiter unten. Aber hier geht es ja gerade um die Mathematik hinter solchen Tests ...)

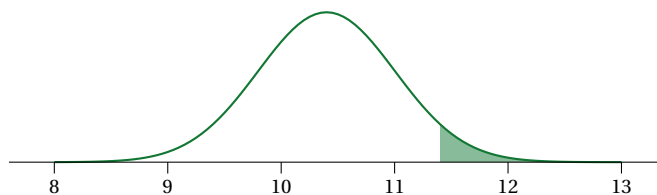
Wenn Ihr Modell richtig ist, dann kann man jedes bisherige Klausurergebnis als **Realisierung einer Zufallsvariable** mit der oben genannten Verteilung betrachten. Wenn die Nullhypothese stimmt, dann sind die 20 Ergebnisse Ihrer Studie Realisierungen unabhängiger Zufallsvariablen X_1 bis X_{20} mit ebendieser Verteilung und die Durchschnittsnote $\overline{X}_{20} = \frac{1}{n} \sum_{i=1}^{20} X_i$ ist normalverteilt mit Erwartungswert 10.4 und Standardabweichung $2.7/\sqrt{20} \approx 0.6$ – oder der Erwartungswert ist sogar noch kleiner. (Dabei nutzen wir aus, dass die Summe von unabhängigen normalverteilten Zufallsvariablen wieder normalverteilt ist. Das ist eine besondere Eigenschaft der Normalverteilung, die andere Verteilungen in der Regel *nicht* haben.)

Aufgabe 39.1. Die Alternativhypothese besagt also, dass die 20 Klausurnoten normalverteilt sind mit Standardabweichung $\sigma = 2.7$ und einem größeren Erwartungswert μ' als 10.4. In der Einleitung hatte ich gesagt, dass man in der Statistik versucht, mithilfe von Daten die Verteilung herauszubekommen. Ist es dann nicht ein Widerspruch, wenn wir die Verteilung im Modell schon voraussetzen?

Das sogenannte *Signifikanzniveau* aus Schritt (H3) legen typischerweise nicht Sie selbst fest. Je nach wissenschaftlicher Disziplin wird es Richtwerte geben, an die Sie sich halten müssen, wenn Sie Ihre Ergebnisse *publizieren* wollen. Wir nehmen für unser Beispiel den in vielen Bereichen verbreiteten Wert von fünf Prozent. Erwähnenswert ist an dieser Stelle jedoch, dass das Signifikanzniveau eine willkürliche Konvention ist, für die es keine mathematische Begründung gibt. Statt 5 hätte man ebenso 10, 1 oder 4.2 nehmen können.

Signifikanzniveau

Testplanung und Signifikanzniveau zusammen ergeben die Situation, die in der folgenden Grafik dargestellt ist.



Wir sehen einen Ausschnitt der *Dichtefunktion* von $\overline{X}_{20}$ für den Fall, dass der Hagebuttentee nicht wirkt. Der grün hervorgehobene Bereich verläuft horizontal von circa 11.39 bis ∞ und ist so gewählt, dass er fünf Prozent der Gesamtfläche unter der Kurve einnimmt.

Aufgabe 39.2. Verständnisfrage zur Grafik: Wie groß ist die Wahrscheinlichkeit, dass $\overline{X}_{20}$ einen Wert annimmt, der kleiner als 11.39 ist? Wie groß ist die Wahrscheinlichkeit, dass genau der Wert 11.39 angenommen wird?

Aufgabe 39.3. Bei der Klausur fehlen zwei Studierende, weil Sie den Hagebuttentee nicht vertragen haben. Sie müssen die Studie also mit 18 statt mit 20 Teilnehmern durchführen. Wie ändert sich die Kurve dadurch? Fängt der grüne Bereich dann weiter rechts oder weiter links an?

Schließlich kommt der große Tag. Die Klausur wird geschrieben, Sie sitzen in Ihrem Büro und berechnen die Durchschnittsnote. Das Ergebnis ist 11.43. Das war Schritt (H4).

Aufgabe 39.4. Wieso ist es offensichtlich, dass Sie sich verrechnet haben?

Damit sind wir bei Schritt (H5). Was ist nun bei der Sache herausgekommen? Weil die Zahl 11.43 im grünen Bereich rechts von 11.39 liegt, haben Sie einen *statistisch signifikanten* Beleg dafür, dass Hagebuttentee sich tatsächlich positiv auf die Mathezensur auswirkt!

Bevor wir darüber nachdenken, was das bedeutet, sei noch einmal wiederholt, dass Sie sich als Anwender typischerweise einen Großteil der mathematischen Arbeit, die gerade vorgeführt wurde, sparen. Nachdem Sie sich für einen Test entschieden haben, werden Sie den Rest einen Computer erledigen lassen. Eine Statistiksoftware wie *R* oder eine Bibliothek wie *SciPy* für PYTHON kennt die gängigen Tests. Sie müssen lediglich Ihre Daten eingeben und bekommen die Ergebnisse dann auf dem Silbertablett serviert.

Aufgabe 39.5. Bei einem Gauß-Test geht es immer um Normalverteilungen. Sie müssen eigentlich nur vier Zahlen eingeben. Im Beispiel wären drei davon 10.4 und 2.7 für Erwartungswert und Standardabweichung der Nullhypothese und 11.43 für das Ergebnis Ihrer Studie. Was wird wohl die vierte Zahl sein?



Nun zur Interpretation der ganzen Geschichte. Das ist aus meiner Sicht der wichtigste Aspekt, weil ich befürchte, dass viele Menschen solche Tests einsetzen, ohne zu verstehen, was sie eigentlich tun.

- Zunächst einmal sollte man sich klarmachen, was dem Zufall unterworfen ist und was nicht. Das zufällige Ereignis ist nämlich Ihre Studie! Wenn wir die Nullhypothese voraussetzen, dann sind die 20 Klausurnoten Realisierungen von Zufallsvariablen. Insbesondere heißt das, dass sich bei einer Wiederholung der Studie höchstwahrscheinlich eine andere Durchschnittsnote ergeben würde.
- Was Sie eigentlich herausfinden wollen, hat jedoch nichts mit Wahrscheinlichkeiten zu tun. Entweder ist Ihre Alternativhypothese wahr oder nicht. Es ergibt keinen Sinn, zu sagen, sie sei zu 78 % wahr.¹⁷
- Aber wie kommt die Aussage des Tests zustande und was bedeutet sie? Sie besagt, dass unter der Annahme der Nullhypothese die Wahrscheinlichkeit für eine Durchschnittsnote wie 11.43 geringer als 5 % ist. Weil das eine sehr geringe Wahrscheinlichkeit ist, scheint es sinnvoll zu sein, statt der Nullhypothese die Alternativhypothese für richtig zu halten.
- Es sei aber noch einmal wiederholt, dass die Zahl 5 % gewissermaßen aus der Luft gegriffen ist. Durch sie wird willkürlich festgelegt, was im vorherigen Absatz unter „sehr gering“ zu verstehen ist: nämlich *per definitionem* alles, was unter 5 % liegt. Und der Begriff *statistisch signifikant* bedeutet nichts anderes als: Die Wahrscheinlichkeit, dieses spezifische Testergebnis (oder ein im Sinne der Alternativhypothese noch besseres) zu erzielen, liegt unterhalb einer vorab festgelegten Zahl.
- Ist 5 % wirklich eine sehr geringe Wahrscheinlichkeit? Das kann man sich folgendermaßen überlegen: Wenn Hagebuttentee keinerlei Einfluss auf die Mathenote hat, Sie aber tapfer Ihr Experiment oft genug wiederholen, dann wird der Test trotzdem im Durchschnitt in etwa 5 % aller Fälle (also bei jedem zwanzigsten Versuch) einen statistisch signifikanten Beleg für die Wirkung von Hagebuttentee behaupten.

statistisch signifikant

Welche Bedeutung haben Null- und Alternativhypothese? Wie bereits erwähnt sollte das, was Sie eigentlich gerne als Ergebnis hätten, die Alternativhypothese sein, die man als H_1 bezeichnet, und nicht die Nullhypothese H_0 . Warum ist das so? Aufgrund der wahrscheinlichkeitstheoretischen Natur des Tests sind Fehler unvermeidbar. Man unterscheidet zwei Sorten von Fehlern:

¹⁷ Man kann jedoch Wahrscheinlichkeit als Maß dafür interpretieren, für wie glaubhaft man eine Aussage hält. Das fällt in den Bereich der sogenannten *bayesschen Statistik*. Siehe dazu das [am Anfang des Kapitels erwähnte Video](#).

	H_0 ist wahr.	H_1 ist wahr.
Test entscheidet sich für H_0 .	OK	Fehler 2. Art
Test entscheidet sich für H_1 .	Fehler 1. Art	OK

Entscheidet sich der Test für H_1 , obwohl H_0 wahr ist, dann spricht man von einem *Fehler 1. Art*. Die Wahrscheinlichkeit dafür, so einen Fehler zu machen, kann man begrenzen, nämlich durch das Signifikanzniveau. Von einem *Fehler 2. Art* spricht man hingegen, wenn der Test sich umgekehrt für H_0 entscheidet, obwohl H_1 wahr ist. Diese Wahrscheinlichkeit kennt man im Allgemeinen bei Tests *nicht*.

Fehler 1. Art
Fehler 2. Art

Das ist der Grund für die Asymmetrie von Nullhypothese H_0 und Alternativhypothese H_1 . Man wählt sie so, dass der Fehler 1. Art „schlimmer“ als der Fehler 2. Art ist, weil man die Wahrscheinlichkeit für den „schlimmen“ Fehler kontrollieren kann. Das ist so wie bei einem Gerichtsverfahren, bei dem das Prinzip *in dubio pro reo* gilt. Man hält den Fehler, einen Unschuldigen zu verurteilen, für schlimmer als den, einen Schuldigen freizusprechen. Daher soll es nur dann zu einem Urteilsspruch kommen, wenn die Indizien für die Schuld des Angeklagten deutlich überwiegen. Wenn Sie einen statistischen Test verwenden, sollten Sie sich bei der Wahl der Hypothesen daher in die Rolle der Staatsanwaltschaft versetzen. Wählen Sie als *Alternativhypothese* das, was Sie „beweisen“ wollen.

Was ist der p -Wert? Eigentlich geht es bei Hypothesentests nur um eine Ja-Nein-Entscheidung: ob man die Nullhypothese zugunsten der Alternativhypothese verwerfen soll oder nicht. Ihre Software wird Ihnen aber typischerweise noch eine Zahl auswerfen, die man den p -Wert nennt; im konkreten Beispiel ungefähr 0.044. Diese Zahl wird oft als Wahrscheinlichkeit missverstanden, aber sie ist keine! Sie ist das minimale Signifikanzniveau, für das Ihr Ergebnis gerade noch signifikant gewesen wäre. (Und sie ist selbst Realisierung einer Zufallsvariablen. Jedes Experiment liefert einen anderen p -Wert.)

p -Wert

Liegt der p -Wert wie in diesem Fall sehr dicht am Signifikanzniveau, so ist das ein Indiz dafür, dass das Ergebnis recht „wackelig“ ist. Im konkreten Beispiel wäre eine etwas schlechtere Durchschnittsnote von 11.39 schon nicht mehr ausreichend gewesen. Und der obige Schnitt von 11.43 hätte bei 18 statt 20 Teilnehmern (siehe Aufgabe 39.3) ebenfalls nicht mehr für ein signifikantes Resultat gereicht.

Welchen Test soll ich verwenden? Würde sich etwas ändern, wenn die Alternativhypothese einfach wäre, dass der Hagebuttentee die Zensuren ändert, sie aber durch ihn evtl. auch schlechter werden können? Ja, man würde dann von einem *einseitigen* zu einem *zweiseitigen* Gauß-Test wechseln und andere Zahlen erhalten. (Siehe auch Aufgabe 39.7.) Den Gauß-Test habe ich oben allerdings nur ausgewählt, weil man ihn leicht erklären kann. In den meisten Fällen passt das hinter ihm stehende Modell nicht zur realen Testsituation. Kennt man beispielsweise die Standardabweichung nicht (was typischerweise der Fall sein wird), so wendet man eher den *t-Test* an. Doch das ist noch lange nicht alles. Wir haben den Mittelwert einer *Stichprobe*¹⁸ mit einem (angeblich) bekannten Mittelwert verglichen. Man kann jedoch auch *zwei Stichproben* mit unbekannten Mittelwerten vergleichen.

¹⁸Das waren die 20 Klausurergebnisse.

Könnten sich in beiden auch die Standardabweichungen unterscheiden, sollte man statt des t -Tests den *Welch-Test* verwenden. Alle bisher erwähnten Tests setzen allerdings voraus, dass die Daten normalverteilt sind. Was macht man, wenn das nicht der Fall ist? Auch dafür gibt es Tests. Es gibt ein solches *Sammelsurium von Tests*, dass es sinnlos wäre, die hier alle aufzuführen. Der Sinn des Abschnitts war, die Grundprinzipien solcher Tests zu vermitteln. Wenn Sie wirklich mal einen benötigen, müssen Sie ein Statistik-Lehrbuch konsultieren oder im Internet nach einen „Entscheidungsbaum“ für statistische Tests suchen.

Die folgenden Aufgaben spielen noch einmal den Mechanismus eines Hypothesentests durch. Dabei wird ein sehr einfaches Beispiel verwendet, für das man keine speziellen Tests nachschlagen muss.

Aufgabe 39.6. Sie vermuten, dass eine Münze (mit den Seiten 0 und 1) manipuliert ist und zu oft die Seite 1 erscheint. Wenn Sie das mit einem Hypothesentest überprüfen wollen, was sind dann Null- und Alternativhypothese?

Aufgabe 39.7. Wie in Aufgabe 39.6 haben Sie es mit einer vermutlich manipulierten Münze zu tun. Sie haben zwar diesen Verdacht, aber Sie haben bisher keinen Anhaltspunkt dafür, welche Seite (0 oder 1) die bevorzugte ist. Wie würden nun Null- und Alternativhypothese aussehen?

★**Aufgabe 39.8.** Wie groß ist bei einer fairen Münze die Wahrscheinlichkeit dafür, dass bei 20 Würfeln 15-mal die Seite 0 oben liegt und 5-mal die Seite 1?

Aufgabe 39.9. Sie führen den in Aufgabe 39.7 beschriebenen Hypothesentest durch, indem Sie die Münze 20-mal werfen. Für die *Zufallsvariable* X , die die Anzahl der Würfe zählt, bei denen die Seite 1 oben liegt, gibt die folgende Tabelle die *gerundeten* Wahrscheinlichkeiten für einen fairen Würfel an.

k	0	1	2	3	4	5	6
$P(X = k)$	0.0000	0.0000	0.0002	0.0011	0.0046	0.0148	0.0370
k	7	8	9	10	11	12	13
$P(X = k)$	0.0739	0.1201	0.1602	0.1762	0.1602	0.1201	0.0739
k	14	15	16	17	18	19	20
$P(X = k)$	0.0370	0.0148	0.0046	0.0011	0.0002	0.0000	0.0000

Das Signifikanzniveau sei 5 %. Entscheiden Sie auf der Basis dieser Zahlen, bei welchen Ergebnissen man die Nullhypothese verwerfen und damit den Würfel als wahrscheinlich gefälscht beurteilen würde.

Aufgabe 39.10. Was würde sich bei Aufgabe 39.9 qualitativ ändern, wenn man die Anzahl der Würfe erhöht und ansonsten nichts ändert? Würden Fehler erster Art unwahrscheinlicher werden?

Aufgabe 39.11. Warum kann man in Aufgabe 39.9 – wie bei den meisten Hypothesentests – die Wahrscheinlichkeit für einen Fehler zweiter Art nicht angeben?

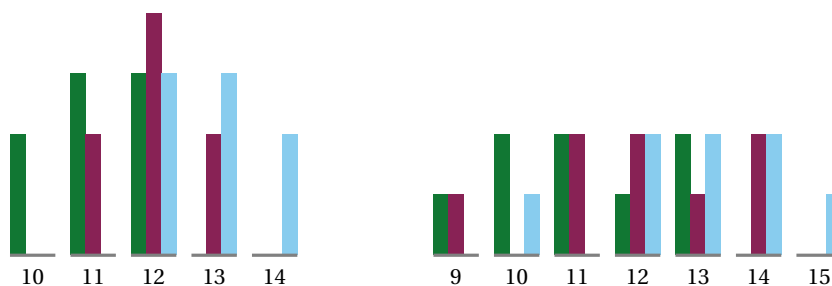
Zum Schluss noch ein kurzes Beispiel für eine andere Klasse von Tests. Auch hier handelt es sich um Hypothesentests, aber der Ansatz ist ein etwas anderer. Das

Beispiel – eine Fortsetzung der Hagebuttentee-Story – zeigt nur die einfachste Variante und die Zahlen wurden denkbar einfach ausgewählt, um das prinzipielle Konzept zu illustrieren:

Nachdem der vorherige Test Sie überzeugt hat, dass Hagebuttentee die gewünschte Wirkung hat, wollen Sie wissen, welchen Einfluss die Marke ausübt. Sie lassen daher erneut eine Klausur mit 24 Studierenden schreiben. Diese werden in drei Gruppen à acht Personen aufgeteilt und bekommen am Abend vor der Klausur unterschiedliche Sorten Tee verabreicht, nämlich solchen von den Firmen Grün, Rot und Blau. Weil es sich diesmal um eine Klausur auf höherem Niveau handelt, gehen Sie nicht von bekannten Werten für Erwartungswert und Standardabweichung aus. Voraussetzung ist aber nach wie vor, dass wir es mit einer Normalverteilung zu tun haben. Die (gerundeten) Durchschnittsnoten zeigt die folgende Grafik:

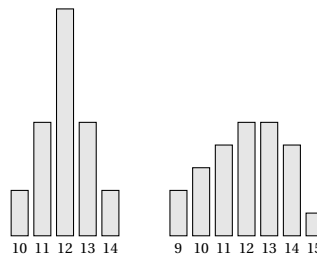


Kann man daraus bereits schließen, dass Firma Blau den besten "Mathe-Tee" herstellt oder könnte das ein zufälliges Ergebnis sein, das zu falschen Schlüssen verleitet? Zunächst einmal sollte inzwischen klar geworden sein, dass bei *jedem* Hypothesentest zufällige Ereignisse zu falschen Resultaten führen können. Das liegt in der Natur der Sache. Es geht lediglich darum, die Wahrscheinlichkeit für bestimmte Fehler nicht zu groß werden zu lassen. Im konkreten Beispiel ist es so, dass die gezeigten Durchschnittsnoten auf sehr unterschiedliche Art zustande gekommen sein können. Dafür zeigt die folgende Grafik zwei Beispiele.



Links liegen die Klausurergebnisse der einzelnen Gruppen jeweils eng beieinander, während rechts ein ziemlicher Kuddelmuddel herrscht. Um den zu entwirren, wird ein Verfahren eingesetzt, dass man *Varianzanalyse* (engl. *Analysis of Variance*, kurz ANOVA) nennt. Im einfachsten Fall, den wir hier vorliegen haben, ist die Nullhypothese, dass alle drei Gruppen denselben Erwartungswert haben. Dann vergleicht man die *Varianz* innerhalb der Gruppen mit der Gesamtvarianz, die man der folgenden Grafik entnehmen könnte. (Wie üblich rechnen Sie das nicht selbst aus, sondern lassen einen Computer die Arbeit machen.)

Varianzanalyse



Sind die Gruppenvarianzen relativ zur Gesamtvarianz klein, so wird die Nullhypothese verworfen. Anderenfalls muss man davon ausgehen, dass die Unterschiede innerhalb der Gruppen auf die insgesamt große Streuung über alle Zufallsvariablen zurückzuführen ist. Im konkreten Fall würde man links einen *p*-Wert von circa 0.1 % erhalten und rechts etwa 10.7 %. Das Ergebnis rechts wäre also bei einem Signifikanzniveau von 5 % nicht zu gebrauchen, während man links guten Gewissens von echten Unterschieden zwischen den verschiedenen Teesorten ausgehen könnte.

Aber Achtung! Bis zu diesem Punkt liefert die Varianzanalyse nur die Aussage, dass die Annahme von Unterschieden sinnvoll ist. Sie sagt nichts darüber aus, ob beispielsweise Tee von der Marke *Blau* besser als Tee von der Marke *Rot* ist! Dafür müssen Sie einen weiteren Test durchführen und ggf. erneut ein Lehrbuch aufschlagen ...

★ Hypothesentests im Computer

Es sei erneut auf die ausführliche Darstellung in *Konkrete Mathematik (nicht nur) für Informatiker* hingewiesen. Hier zeige ich nur zwei kurze Beispiele um anzudeuten, wie einfach Hypothesentests mit Computerhilfe sind – wenn man weiß, was man tut ...

Zuerst das Beispiel mit dem Hagebuttentee. Wie schon erwähnt ist der Gauß-Test selten anwendbar, weil man die Standardabweichung nicht kennen wird. Verwendet man stattdessen den üblicheren *t*-Test, dann könnte das so aussehen:

```
# 20 Klausurergebnisse
L=[11,12,8,15,12,14,11,14,12,14,8,8,12,10,11,14,11,7,13,11]

from scipy.stats import ttest_1samp
ttest_1samp(L,10.4,alternative="greater")
```

Hier ist 10.4 die Durchschnittsnote der Nullhypothese und der Parameter *greater* beschreibt die Alternativhypothese, dass die Werte in der Liste L im Schnitt größer als 10.4 sind.¹⁹

Schließlich noch die rechte Hälfte der Varianzanalyse mit den drei Teemarken:

¹⁹Die Durchschnittsnote ist 11.4 und *nicht* 11.43. Siehe Aufgabe 39.4.

```
green = [11,13,10,9,11,12,13,10]
red = [14,11,12,11,13,12,9,14]
blue = [12,14,10,13,15,14,12,13]

from scipy.stats import f_oneway
f_oneway(green,red,blue)
```


Lösungen zu ausgewählten Aufgaben

Lösung 1.1. (i) ist keine Aussage, sondern ein Ausdruck, der für eine Zahl steht. (iv) ist keine Aussage, weil nicht klar ist, wofür x steht. Die anderen Ausdrücke sind Aussagen, wobei nur (ii) wahr ist.

Lösung 1.2. Man muss dafür lediglich in der Tabelle nachschauen. In tabellarischer Form sieht die Antwort so aus:

0	1	0
1	0	0

Lösung 1.3. Nach Definition hängt beispielsweise der Wahrheitswert von $(A \wedge B)$ nur von den Wahrheitswerten von A und B ab. Für den Wahrheitswert von A gibt es nur die beiden Möglichkeiten 0 und 1, für den von B auch nur. Daher können offenbar nur die vier Kombinationen auftreten, die in der Tabelle gezeigt werden.

Lösung 1.4. Ein Junktor kann nur durch die Anwendung einer Regel in eine Formel gelangen, die gleichzeitig ein Klammersymbol hinzufügt. $(E \wedge \wedge E)$, $(A \wedge B \wedge C)$ und $A \wedge B$ haben daher zu wenige Klammern und können keine Formeln sein. $((A \wedge C))$ hat zu viele Klammern und ist ebenfalls keine Formel. (EE) ist keine Formel, weil in atomaren Formeln immer nur ein Buchstabe vorkommt und alle anderen Formeln offenbar mindestens einen Junktor enthalten. $(A + B)$ ist keine Formel, weil das Symbol $+$ vorkommt, das in der Definition gar nicht erwähnt wird. Lediglich E und $(A \wedge (B \wedge C))$ sind nach der Definition Formeln. Letztere ist durch zweimalige Anwendung der Regel für $(\alpha \wedge \beta)$ entstanden.

Lösung 1.5. Für *eine* Variable gibt es offenbar nur die beiden Möglichkeiten 0 und 1. Fügt man eine weitere Variable hinzu, so kann man jede Belegung der ersten Variable mit einer der beiden möglichen Belegungen für die zweite kombinieren, so dass man insgesamt vier Belegungen erhält. Mit jeder weiteren Variable kommt es offenbar erneut zu einer Verdoppelung, so dass die erste Antwort 8 ist und die zweite 2^n . (Siehe auch Aufgabe 1.3.)

Lösung 1.6. Zunächst noch einmal zur Übung eine ausführliche Wahrheitstafel.

A	B	$(\neg A)$	$(\neg B)$	$(A \wedge (\neg B))$	$(A \Rightarrow B)$	$((\neg A) \vee B)$	$(\neg(A \wedge (\neg B)))$
0	0	1	1	0	1	1	1
0	1	1	0	0	1	1	1
1	0	0	1	1	0	0	0
1	1	0	0	0	1	1	1

Entscheidend ist, dass in den drei letzten Spalten dasselbe steht. Man kann also die Aussage „Aus A folgt B “ auch anders interpretieren und dabei auf die eventuell missverständliche Erwähnung der Folgerung verzichten: „ A ist falsch oder B ist wahr“ (vorletzte Spalte) bzw. „ A und $\neg B$ können nicht beide wahr sein“ (letzte Spalte).

Lösung 1.8. Mit allen Klammern hat die Formel die Form

$$(((A \vee B) \Rightarrow ((\neg A) \vee (C \wedge B))) \Leftrightarrow (A \vee C)).$$

Lösung 1.9. Ihr Ergebnis sollte so aussehen:

$$C \wedge B \vee A \wedge B \Rightarrow (A \vee B) \wedge (B \vee C).$$

Beachten Sie, dass die verbleibenden Klammern nicht entfernt werden dürfen, weil $\wedge$ eine höhere Priorität als $\vee$ hat und man die Formel ganz ohne Klammern als

$$C \wedge B \vee A \wedge B \Rightarrow A \vee (B \wedge B) \vee C$$

interpretieren müsste.

Lösung 1.10. Weil die beiden Formeln nicht äquivalent sind. Verifizieren Sie das mit Wahrheitstafeln, falls es Ihnen nicht ohnehin klar war.

Lösung 1.14. Belegt man A und B beide mit 0 oder beide mit 1, dann ist die Formel wahr. Also ist sie erfüllbar. Haben A und B jedoch unterschiedliche Wahrheitswerte, dann ist die Formel nicht wahr, also ist es keine Tautologie.

Lösung 1.15. $A \vee B$ ist äquivalent zu $\neg(\neg A \wedge \neg B)$. $A \Leftrightarrow B$ ist nach Aufgabe 1.7 äquivalent zu $(A \Rightarrow B) \wedge (B \Rightarrow A)$ und nach Aufgabe 1.6 ist diese Formel äquivalent zu einer, in der nur noch $\wedge$ und $\neg$ vorkommen.

Auf $\vee$, $\Rightarrow$ und $\Leftrightarrow$ könnte man also im Prinzip auch verzichten und Formeln, in denen diese Zeichen verwendet werden, als Abkürzungen für Formeln interpretieren, in denen nur $\wedge$ und $\neg$ vorkommen. Ebenso kann man alles andere (auch $\wedge$) nur mit $\vee$ und $\neg$ ausdrücken. (Vielleicht probieren Sie das mal aus?)

In den Informatik-Vorlesungen lernt man, dass man sogar mit einem einzigen Junktoren auskommen kann, nämlich mit $\overline{\vee}$ (NOR) oder mit $\overline{\wedge}$ (NAND).

Lösung 1.16. Ja, es handelt sich um eine Tautologie (und damit natürlich „erst recht“ um eine erfüllbare Formel). Ihre Wahrheitstafel hatte hoffentlich acht Zeilen – siehe Aufgabe 1.5.

Lösung 1.17. Mit einer Wahrheitstafel kann man nachprüfen, dass (i) und (iii) zum Erfolg führen, (ii) aber nicht.

Lösung 1.19. Setzt man A für „Der Hahn kräht auf dem Mist“ und B für „Das Wetter ändert sich“, dann handelt es sich um die Aussage $A \Rightarrow B \vee \neg B$. Weil das eine Tautologie ist, ist das Sprichwort auf jeden Fall wahr – unabhängig davon, wie das Wetter am nächsten Tag ist.

Lösung 1.20. Man könnte das Rätsel folgendermaßen lösen: A sei die Aussage, dass sich in der ersten Kiste eine Orange befindet, und B die entsprechende Aussage für die zweite Kiste. Das Schild auf der ersten Kiste besagt somit $A \wedge \neg B$ und das auf der zweiten $(A \wedge \neg B) \vee (\neg A \wedge B)$. Das liefert die folgende Tabelle:

A	B	$\neg A$	$\neg B$	$A \wedge \neg B$	$(A \wedge \neg B) \vee (\neg A \wedge B)$
0	0	1	1	0	0
0	1	1	0	0	1
1	0	0	1	1	1
1	1	0	0	0	0

Nur in der zweiten Zeile sind die Wahrheitswerte für die beiden Schilder unterschiedlich. Die muss es sein und deshalb ist B wahr: Sie finden in der zweiten Kiste eine Orange.

Es ist nicht schlimm, wenn Sie das anders gelöst haben, solange Sie zum selben Ergebnis gekommen sind. (Siehe Aufgabe 1.21.) Es sei jedoch auf einen typischen Fehler hingewiesen: Oft kann man beobachten, dass beispielsweise A für „Orange in der ersten Kiste“, B für „Apfel in der ersten Kiste“, C für „Orange in der zweiten Kiste“ und D für „Apfel in der zweiten Kiste“ gesetzt wird. Das ist nicht zielführend, weil dann z. B. A und B nicht unabhängig voneinander sind. Es ist beispielsweise nicht möglich, dass sowohl A als auch B wahr sind, weil sich nach der Aufgabenstellung in jeder Kiste nur ein Stück Obst befindet. B und D sind redundant, weil man sie ohne Informationsverlust durch $\neg A$ bzw. $\neg C$ ersetzen kann.

Lösung 1.21. Zum Beispiel mit $\neg(A \Leftrightarrow B)$ oder $\neg(A \wedge B) \wedge \neg(\neg A \wedge \neg B)$. Es gibt immer viele Möglichkeiten, ein und dieselbe Aussage formal zu repräsentieren. (Sogar immer unendlich viele. Ist Ihnen klar, wieso das so ist?)

Lösung 1.22. Setzen wir A dafür, dass der HSV die Champions League gewinnt, und B dafür, dass ich in die Alster springe. Dann besagen die vier Sätze in dieser Reihenfolge $A \Rightarrow B$, $\neg A \Rightarrow \neg B$, $\neg B \Rightarrow \neg A$ und $A \Rightarrow \neg B$. Mit einer Wahrheitstafel überzeugt man sich, dass nur die erste und die dritte Aussage äquivalent sind.

Lösung 1.23. Sei A die Aussage, dass auf der Buchstabenseite ein Vokal steht, und B die Aussage, dass auf der Zahlenseite eine gerade Zahl steht. Die zweite Regel ist also einfach $A \Rightarrow B$. Diese Regel wird nur dann verletzt, wenn A wahr ist und B falsch. Darum müssen Sie sich die Karte mit dem Buchstaben U (A ist wahr) und die Karte mit der Zahl 7 (B ist falsch) genauer anschauen. Die anderen beiden

Karten müssen Sie nicht umdrehen, denn die Implikation ist unabhängig davon, was auf deren Rückseite gedruckt wurde, immer wahr.

Wie bei Aufgabe 1.20 sollte man auch hier darauf achten, dass man keine unnötigen Aussagen verwendet. Es wäre beispielsweise kontraproduktiv, noch eine Aussage C für „Konsonant auf der Buchstabenseite“ einzuführen, weil ein Buchstabe nicht gleichzeitig Vokal und Konsonant sein kann. C lässt sich durch $\neg A$ ausdrücken.

Lösung 1.24. Bezeichnet man die Aussage „Gauß ist schuldig“ zur Abkürzung mit G und die entsprechenden Aussage für Euler und Hilbert mit E und H , so lassen sich die Ergebnisse der bisherigen Ermittlungen wie folgt zusammenfassen: $(H \Rightarrow E) \wedge (G \vee H \Rightarrow \neg E) \wedge (\neg E \vee \neg H \Rightarrow G)$. Mit einer Wahrheitstafel kann man nun nachprüfen, dass diese Formel nur dann wahr wird, wenn G wahr ist und H und E falsch sind. Daher muss Gauß der (einzige) Täter sein.

Lösung 1.25. Da diese Frage im Abschnitt über Aussagenlogik vorkommt, lösen wir sie natürlich mit Mitteln der Aussagenlogik. Seien A und B die Aussagen, dass Andrew bzw. Bert blau sind. Andrew wird auf die erste Frage mit *ja* antworten, wenn beide blau sind oder wenn er selbst grün ist, wenn also $(A \wedge B) \vee \neg A$ gilt. (Diese Aussage kann man zu $\neg A \vee B$ vereinfachen, obwohl das nicht unbedingt nötig ist.) Die zweite Frage wird er mit *ja* beantworten, wenn beide blau sind oder wenn er grün und Bert blau ist. Das ist also dann der Fall, wenn $(A \wedge B) \vee (\neg A \wedge B)$ gilt. (Man sieht hoffentlich leicht, dass das äquivalent zu B ist, aber auch diese Vereinfachung ist für das Lösen der Aufgabe nicht zwingend notwendig.) Es ergibt sich die folgende Wahrheitstafel:

A	B	$(A \wedge B) \vee \neg A$	$(A \wedge B) \vee (\neg A \wedge B)$
0	0	1	0
0	1	1	1
1	0	0	0
1	1	1	1

In der dritten bzw. vierten Spalte steht jeweils genau dann eine Eins, wenn Andrew auf die erste bzw. zweite Frage mit *ja* geantwortet hat. Da die Astronautin nach der ersten Frage noch nicht die Farben weiß, kann Andrews Antwort nicht *nein* gewesen sein, denn daraus hätte sie sofort schließen können, dass Andrew blau und Bert grün ist. Also können wir diesen Fall vergessen. (Darum ist die dritte Zeile ausgegraut.) Da die Astronautin nach der zweiten Antwort die Farben weiß, muss die zweite Antwort *nein* gewesen sein, denn die Antwort *ja* hätte ihr die Farbe von Andrew nicht verraten. Also sind beide grün.

Lösung 2.1. Wir stellen das tabellarisch dar, wobei ein Kreuz bedeuten soll, dass die Aussage wahr ist:

	X	Y	Z
X	×		
Y	×	×	
Z	×		×

Solche Tabellen sind immer „von links nach oben“ zu lesen. Der zweite Eintrag in der ersten Zeile steht also beispielsweise für die (falsche) Aussage $X \subseteq Y$ und *nicht* für $Y \subseteq X$.

Beachten Sie, dass weder $Y \subseteq Z$ noch $Z \subseteq Y$ gilt. Lemma 2.1 zeigt zwar, dass es viele Ähnlichkeiten zwischen $\subseteq$ bei Mengen und $\leq$ bei Zahlen gibt, aber an dieser Stelle unterscheiden sich die beiden: Wenn man zwei (reelle) Zahlen y und z hat, dann ist immer mindestens eine der beiden Aussagen $y \leq z$ und $z \leq y$ wahr. Für Mengen und $\subseteq$ gilt das offenbar nicht.

Lösung 2.2. Keine der Aussagen ist wahr. Es geht hier lediglich um das Ausräumen des Missverständnisses, dass die Zahl null und die leere Menge etwas miteinander zu tun hätten, weil die Zeichen ähnlich aussehen. Die erste Aussage ist falsch, weil 0 eine Zahl ist und $\emptyset$ eine Menge.²⁰ Die zweite Aussage ist falsch, weil das Singleton $\{0\}$ ein Element enthält (nämlich die Zahl null), $\emptyset$ aber *kein* Element enthält. Die dritte Aussage ist ebenfalls falsch, weil $\emptyset$ keine Elemente hat. Die einzige wahre Aussage, die in diesem Zusammenhang Sinn ergibt, ist $0 \in \{0\}$. (Dass $0 = \{0\}$ *nicht* wahr ist, hatten wir schon.)

Lösung 2.3. Die Lösungen sind in der Reihenfolge der Aufgabenstellung $\{3, 7\}$, $\{7\}$, $\{1, 2, 3, 5, 6, 7\}$, $\{1, 2, 3, 4, 5, 6, 7\}$, $\{3, 4, 5, 6, 7\}$, $\{5, 6, 7\}$, $\{1, 5\}$, $\{1\}$ und $\{1, 5, 7\}$.

Wenn Sie noch mehr Übung brauchen, dann können Sie sich selbst Aufgaben stellen und Ihre Lösungen überprüfen, indem Sie diese in [Wolfram Alpha](#)²¹ eingeben. Für den Ausdruck $\{1, 2\} \cup \{3, 4, 5\}$ kann man z. B.

union of $\{1, 2\}$ and $\{3, 4, 5\}$

eingeben. Für den Durchschnitt verwendet man *intersection* und für die Differenz *set difference*.

Lösung 2.5. $A \triangle B = \{5, 9, 12\}$. Grundsätzlich gilt immer $A \triangle B = B \triangle A$, was aus Lemma 2.2 folgt. Überzeugen Sie sich davon ggf. mit einem Mengendiagramm.

Lösung 2.6. Mithilfe eines Mengendiagramms kann man sich klarmachen, dass $A \triangle B = (A \setminus B) \cup (B \setminus A)$ gilt.

Lösung 2.7. Betrachten Sie beispielsweise die Mengen

$$A = \{1, 3, 5, 7, 9, 11, 13, 15\},$$

$$B = \{2, 3, 6, 7, 10, 11, 14, 15\},$$

$$C = \{4, 5, 6, 7, 12, 13, 14, 15\} \text{ und}$$

$$D = \{8, 9, 10, 11, 12, 13, 14, 15\}$$

und versuchen Sie, die Zahlen von 1 bis 15 an den korrekten Stellen im Diagramm zu platzieren. Die Zahl 5 ist Element von A und C , aber weder von B noch von D . Für die gibt es beispielsweise keinen angemessenen Platz.

²⁰In der Informatik würde man sagen, dass es sich um unterschiedliche **Typen** von Objekten handelt.

²¹Achtung: Schauen Sie bei Wolfram Alpha immer im Bereich *Input Interpretation* nach, ob das System auch das macht, was Sie wollen!

Lösung 2.8. Das könnte man z. B. so machen, wobei die Zahlen aus der Lösung von Aufgabe 2.7 gleich eingetragen wurden.

C	4	5	7	6
	12	13	15	14
D	8	9	11	10
	1	3	2	
A				B

Auch für mehr als vier Mengen ist das möglich. Beispiele dafür finden Sie bei Interesse in meinem Video [7 Vielecke](#).

Lösung 3.1. Das war hoffentlich nur eine kleine Fingerübung für Sie. Gemeint war $x < y \vee x = y$.

Lösung 3.2. In diesem Fall war $x \leq y \wedge \neg(x = y)$ bzw. $x \leq y \wedge x \neq y$ gemeint.

Lösung 3.3. Man kann 0 oder 1 einsetzen. Weitere Möglichkeiten gibt es nicht. Mit anderen Worten: Setzt man für x irgendeine andere Zahl ein, dann wird aus $x \cdot x = x$ eine falsche Aussage.

Lösung 3.4. Die einfachsten Beispiele wären $x = x$ für $A(x)$ und $x \neq x$ für $B(x)$. Aber man kann natürlich auch etwas kreativer sein und z. B. $2x = x + x$ für $A(x)$ wählen oder $x = x + 1$ für $B(x)$.

Lösung 3.5. Es gibt mehr als eine Möglichkeit, das zu machen. Hier zwei Vorschläge:

$$(A_4(n) \wedge \neg A_{100}(n)) \vee A_{400}(n)$$

$$A_4(n) \wedge (A_{400}(n) \vee \neg A_{100}(n))$$

Lösung 3.6. $\exists n \in M A(n)$ ist wahr, weil man für n eine der Zahlen 1 oder 2 einsetzen kann. $\exists n \in M B(n)$ ist wahr, weil man für n eine der Zahlen 4 oder 5 einsetzen kann. $\exists n \in M (A(n) \wedge B(n))$ ist nicht wahr, denn für *keine* Zahl m („erst recht nicht“ für eine aus der Menge M) ist $m < 3 \wedge m > 3$ wahr.

Beachten Sie den feinen Unterschied: Die Aussage $(\exists n \in M A(n)) \wedge (\exists n \in M B(n))$ ist als Konjunktion zweier wahrer Aussagen natürlich wahr. Diese Aussage weicht von der vorherigen lediglich durch die Klammerung ab, aber das ist entscheidend.

Lösung 3.7. $\forall n \in M A(n)$ ist nicht wahr, weil z. B. $A[5]$ falsch ist. $\forall n \in M B(n)$ ist nicht wahr, weil z. B. $B[1]$ falsch ist. Daher ist auch $(\forall n \in M A(n)) \vee (\forall n \in M B(n))$ falsch. $\forall n \in M (A(n) \vee B(n))$ ist jedoch wahr, denn für jedes Element m von M ist zumindest eine der beiden Aussagen $A[m]$ oder $B[m]$ wahr.

Lösung 3.8. Die erste Aussage ist wahr. Die zweite Aussage ist falsch, weil $x < y$ falsch wird, wenn man für x und y jeweils die Zahl 3 einsetzt, die in beiden Mengen enthalten ist. Die dritte Aussage ist wahr, weil z. B. $2 < 4$ wahr ist. Die vierte Aussage

ist wahr, weil $3 \leq 3$ wahr ist. Die fünfte Aussage ist falsch, denn man kann keine Elemente m von M und n von N finden, so dass $n < m$ gilt.

Lösung 3.9. Die erste Aussage ist wahr. Man wählt für y einfach dasselbe Element von M aus, das man auch für x eingesetzt hat. Die zweite Aussage ist falsch, denn es gibt nicht *ein* Element von M , das identisch mit *allen* Elementen von M ist. Für die restlichen Aussagen nur eine Zusammenfassung: (iii), (vi), (viii), (ix) und (x) sind wahr, (iv), (v), (vii) und (xi) sind falsch.

Beachten Sie, dass die sechste Aussage beispielsweise nicht mehr wahr wäre, wenn M nur ein Element hätte. Allerdings wäre in diesem Fall die zweite Aussage wahr.

Lösung 3.10. Diese Aufgabe sollte Ihre Aufmerksamkeit darauf lenken, dass man in der Mathematik auch immer genau über „Randfälle“ nachdenken muss. Im Prinzip ist die Folgerung korrekt, allerdings mit einer Einschränkung: Man darf so nicht schließen, wenn M die leere Menge ist. Gilt $M = \emptyset$, dann ist $\forall x \in M A(x)$ unabhängig von $A(x)$ immer wahr, denn sonst müsste es nach Lemma 3.1 ein Element m von M geben, für das $A[m]$ falsch wäre. Aber es gibt ja gar keine Elemente von M . Aus diesem Grund kann auch $\exists x \in M A(x)$ nicht wahr sein.

Lösung 3.11. Ja. Nach Lemma 3.1 kann man $\forall x \in M A(x)$ als Abkürzung für die Aussage $\neg \exists x \in M \neg A(x)$ betrachten. (Umgekehrt könnte man $\exists$ auch durch $\forall$ ersetzen.)

Lösung 3.12. Das wäre so etwas wie

$$\forall x \in M \exists y \in N (y \mid x \wedge \forall z \in N (z \mid x \Rightarrow z = y)).$$

In Worten besagt diese Formel, dass es ein y in N gibt, das x teilt, und dass zudem jedes z aus N , das x teilt, dasselbe Element wie y sein muss.

Lösung 4.1. $B = \{10, 20, 50\}$ und $C = \{20, 40\}$.

Lösung 4.2. Zwei Möglichkeiten wären

$$\mathbb{N}^+ = \{n : n \in \mathbb{N} \wedge n \geq 1\} \quad \text{und} \quad \mathbb{N}^+ = \{n : n \in \mathbb{N} \wedge n \neq 0\}.$$

Ihnen fallen sicher noch andere ein.

Lösung 4.3. Zum Beispiel so:

$$Y = \{n : n \in \mathbb{N}^+ \wedge n \leq 100\}$$

$$Z = \{n : n \in \mathbb{N} \wedge n > 2 \wedge 2 \nmid n\}$$

Lösung 4.4. Am naheliegendsten ist wohl $\{k : k \in \mathbb{N} \wedge 2 \mid k\}$.

Lösung 4.5. Die Definition von $\mathbb{P}$ kann man als

$$\mathbb{P} = \{n : n \in \mathbb{N} \wedge n > 1 \wedge \forall k \in \mathbb{N} (k \mid n \Rightarrow k = 1 \vee k = n)\}$$

übersetzen. Die Menge der zusammengesetzten Zahlen ist $\mathbb{N}^+ \setminus (\mathbb{P} \cup \{1\})$.

Lösung 4.6. $[3]$ ist die Menge $\{1, 2, 3\}$.

$[0]$ ist die leere Menge, denn es gibt keine²² positive natürliche Zahl k mit $k \leq 0$. So etwas werden Sie oft in mathematischen Texten finden. Wenn Sie beispielsweise $x_1, \dots, x_n$ lesen, dann sind damit wahrscheinlich n Variablen gemeint, aber es würde eine Mathematikerin nicht überraschen, dass darin auch der Fall $n = 0$ inbegriffen ist, bei dem es dann eben *keine* Variable gibt.

Falls Sie die Frage nach $[0]$ für unnötig kleinlich halten, dann bedenken Sie bitte, dass diese Menge perfekt zu den anderen passt: $[7]$ hat 7 Elemente, $[42]$ hat 42 Elemente und $[0]$ hat 0 Elemente. Das werden wir ab und zu noch benötigen.

$[\pi]$ würde zwar vielleicht durchaus Sinn ergeben, aber bei dieser Frage ging es darum, dass $[\pi]$ gar nicht definiert ist. Die Definition gilt nur für natürliche Zahlen n . Denken Sie daran, immer aufmerksam zu lesen!

Lösung 4.7. A ist $\{1, 4, 9, 16, 25, 36, 49\}$. Man könnte

$$A = \{k : \exists n \in [7] \ k = n^2\}$$

schreiben. Daran sieht man, dass man im Prinzip auf die erweiterete Schreibweise (4.3) verzichten könnte, sie aber sehr praktisch ist.

Lösung 4.8. $\{2n : n \in \mathbb{N}\}$ und *nicht* $\{2n : n \in \mathbb{N}^+\}$, weil dann die Null fehlen würde.

Lösung 4.9. $\{3n + 1 : n \in \mathbb{N}\}$. Es gibt immer mehr als eine Möglichkeit, aber dies ist wohl die einfachste.

Lösung 4.10. A ist $\{0, 1, 4, 9, 16, 25\}$ und B ist $\{0, 1, 2\}$. Der didaktische Sinn dieser Aufgabe ist natürlich, Sie auf den Unterschied zwischen den beiden ähnlich aussehenden Definitionen aufmerksam zu machen.

Lösung 4.11. $M = \{-2, -1, 0, 1, 2\}$. Sie haben hoffentlich die negativen Zahlen nicht vergessen. (Man vergleiche mit der Menge B in Aufgabe 4.10.)

Lösung 4.12. Die folgenden Darstellungen sind lediglich Vorschläge. Andere Beschreibungen könnten auch korrekt sein.

$$A = \mathbb{N} \cup \{-3, -2, -1\} = \{z : z \in \mathbb{Z} \wedge z > -4\}$$

$$B = \mathbb{Z}$$

$$C = \{n : n \in \mathbb{Z} \wedge n > 3\} = \{n : n \in \mathbb{N} \wedge n \geq 4\}$$

$$D = \mathbb{Z} \setminus \{-3, -2, -1, 0, 1, 2, 3\}$$

$$E = \{n : n \in \mathbb{N} \wedge 3 \nmid n\}$$

Lösung 4.13. Die Antwort ist $\{1, 2, 3, 4, 6, 8, 9, 12, 16\}$. Man erinnere sich, dass die Aussageform $m, n \in [4]$ eine Abkürzung für $m \in [4] \wedge n \in [4]$ ist.

²²Genau lesen! Es gibt sehr wohl eine natürliche Zahl k mit $k \leq 0$, nämlich $k = 0$. Aber in der Definition steht $\mathbb{N}^+$, es geht also nur um *positive* natürliche Zahlen.

Lösung 4.14. Die folgenden Darstellungen sind lediglich Vorschläge. Andere Beschreibungen könnten auch korrekt sein.

$$A = \{q : q \in \mathbb{Q} \wedge q \geq 0\}$$

$$B = \{q : q \in \mathbb{Q} \wedge 0 \leq q < 1\}$$

$$C = \{q : q \in \mathbb{Q} \wedge -1 < q < 1\}$$

$$D = \mathbb{Z}$$

$$E = \mathbb{Q}$$

Lösung 4.15. Das war hoffentlich einfach:

$$\mathbb{N}^+ = \{n \in \mathbb{N} : n \geq 1\}$$

$$[n] = \{k \in \mathbb{N}^+ : k \leq n\}$$

$$X = \{m \in \mathbb{Z} : m < 42\}$$

Lösung 4.16. Für $B = \{x \in \mathbb{N} : x^2 < 6\}$.

Lösung 4.17. $A_0 = \{0\}$, $A_3 = \{3, 4, 5, 6\}$, $A_5 = \{5, 6, 7, 8, 9, 10\}$.

Lösung 4.18. Aussage (i) ist falsch. Sie besagt, dass es eine natürliche Zahl gibt, die nicht gerade, aber durch 4 teilbar ist. Das geht natürlich nicht. Aussage (ii) ist wahr. Sie besagt, dass es eine natürliche Zahl gibt, die gerade, aber nicht durch 4 teilbar ist. Dafür gibt es unendlich viele Beispiele, etwa 2, 6, oder 10. Aussage (iii) ist wahr. Sie besagt, dass alle natürlichen Zahlen durch 4 teilbar sind, *wenn* alle natürlichen Zahlen gerade sind. Da die erste Aussage schon falsch ist (weil eben *nicht* alle natürlichen Zahlen gerade sind), kann auch die zweite ruhig falsch sein.

Aussage (iv) ist falsch. Sie besagt, dass jede gerade natürliche Zahl durch 4 teilbar ist. Gegenbeispiele sind die unter (ii) genannten Zahlen. Aussage (v) ist wahr. Sie besagt, dass jede natürliche Zahl, die durch 4 teilbar ist, auch gerade ist. Aussage (vi) ist falsch. Sie besagt, dass jede natürliche Zahl genau dann gerade ist, wenn sie durch 4 teilbar ist. Wenn das wahr wäre, dann wären sowohl (iv) als auch (v) wahr. Und Aussage (vii) ist auch falsch. Sie besagt, dass jede natürliche Zahl, die nicht durch 4 teilbar ist, auch nicht gerade ist. Man kann hier wieder die Gegenbeispiele aus (ii) verwenden.

Beachten Sie, dass (iii) und (iv) ganz unterschiedliche Aussagen sind, obwohl sie sehr ähnlich aussehen! Für (iii) hätte man auch $(\forall n \in \mathbb{N} \ 2 \mid n) \Rightarrow (\forall m \in \mathbb{N} \ 4 \mid m)$ schreiben können.

Lösung 4.19. Aussage (i) ist wahr. Man kann zu vorgegebenem m z. B. die Zahl $n = m + 1$ wählen. Aussage (ii) ist jedoch falsch, denn hier geht es um eine einzige Zahl $n \in \mathbb{N}$, die größer als *jede* Zahl m sein soll. So etwas gibt es selbstverständlich nicht. (Falls Sie die Idee hatten, man könne $n = \infty$ nehmen, dann sind Sie nicht allein. Es handelt sich allerdings um einen typischen Fehler. ∞ ist keine Zahl!)

(iii) und (iv) sind beide wahr, weil man jeweils die Null für n einsetzen kann.

Lösung 4.20. $A \cap B = \{x : x \in A \wedge x \in B\}$ und $A \setminus B = \{x : x \in A \wedge \neg x \in B\}$ oder auch kürzer $A \setminus B = \{x : x \in A \wedge x \notin B\}$.

Beachten Sie die absichtliche Ähnlichkeit der Symbole $\wedge$ und $\cap$ sowie $\vee$ und $\cup$.

Lösung 4.21. Z. B. $A \triangle B = \{x : (x \in A \vee x \in B) \wedge \neg(x \in A \wedge x \in B)\}$. Es gibt allerdings noch andere Möglichkeiten. Mit dem auf Seite 10 erwähnten exklusiven Oder $\dot{\vee}$ geht es besonders knapp: $A \triangle B = \{x : x \in A \dot{\vee} x \in B\}$.

Lösung 4.22. Für (xix) gilt

$$\begin{aligned} A \setminus (B \cap C) &= \{x : x \in A \wedge \neg(x \in B \wedge x \in C)\} \quad \text{und} \\ (A \setminus B) \cup (A \setminus C) &= \{x : (x \in A \wedge \neg x \in B) \vee (x \in A \wedge \neg x \in C)\} \\ &= \{x : x \in A \wedge (\neg x \in B \vee \neg x \in C)\}. \end{aligned}$$

Nach den de-morganschen Gesetzen sind die beiden die Mengen beschreibenden Aussagen äquivalent. Für (xx) verläuft die Argumentation ebenso, es werden lediglich Konjunktion und Disjunktion vertauscht.

Lösung 4.23. Es gilt $M \subseteq N$: Ist x ein Element von M , so gilt $x \in X$ und $A(x)$ nach Definition von M . Nach Aufgabenstellung folgt daraus $B(x)$ und x ist daher ein Element von N .

Lösung 4.24. Für $\mathbb{P} \subset \mathbb{N}$ kann man eine zusammengesetzte Zahl wie 42 nehmen, für $\mathbb{N} \subset \mathbb{Z}$ eine negative Zahl wie -42 und für $\mathbb{Z} \subset \mathbb{Q}$ einen echten Bruch wie $1/2$.

Ist p irgendeine beliebige ganze Zahl, dann ist $p/1$ ein Element von $\mathbb{Q}$ und es gilt natürlich $p/1 = p$. Das ist die Begründung dafür, dass $\mathbb{Z}$ eine Teilmenge von $\mathbb{Q}$ ist.

Lösung 4.25. Es ergibt sich

$$\begin{aligned} A &= \{n \in \mathbb{N} : n^2 < 400 \wedge 12 \mid n^2\} \quad \text{und} \\ B &= \{n \in \mathbb{N} : n^2 < 400 \wedge 12 \mid n\}. \end{aligned}$$

$C = \{0, 1, 2, \dots, 19\}$ ist die Menge der natürlichen Zahlen, deren Quadrat kleiner als 400 ist. Nur zwei der Elemente von C sind durch 12 teilbar, nämlich 0 und 12. Die Menge der zugehörigen Quadrate ist $\{n^2 : n \in C\} = \{0, 1, 4, 9, \dots, 361\}$ und in dieser Menge sind 0, 36, 144 und 324 durch 12 teilbar. Daher haben wir

$$A = \{0, 6, 12, 18\} \quad \text{und} \quad B = \{0, 12\}.$$

Die beiden Mengen sind nicht gleich.

Diese Aufgabe soll deutlich machen, dass es im Allgemeinen nicht von Vorteil wäre, die von Studierenden oftmals gehassten Formeln durch „normale“ Sprache zu ersetzen. Ich hoffe, Sie stimmen mir zu, dass die beschreibende Mengenschreibweise in diesem Fall wesentlich klarer ist als die Erklärung in der Aufgabenstellung. (Nebenbei gilt das insbesondere auch dann, wenn ein Sachverhalt in einer Sprache formuliert ist, die nicht Ihre Muttersprache ist. Die Formelsprache ist international.)

Lösung 5.1. Das sollte so aussehen:

$$\forall a, b \in \mathbb{N} \quad ab \in \mathbb{N}$$

$$\forall a, b \in \mathbb{N} \quad a + b = b + a$$

$$\forall a, b, c \in \mathbb{N} \quad (a + b) + c = a + (b + c)$$

$$\forall a, b \in \mathbb{N} \quad ab = ba$$

$$\forall a, b, c \in \mathbb{N} \quad (ab)c = a(bc)$$

$$\forall a, b, c \in \mathbb{N} \quad a(b + c) = ab + ac$$

$$\forall n \in \mathbb{N} \quad n + 0 = n$$

$$\forall m, n \in \mathbb{N} \quad (n + m = n \Rightarrow m = 0)$$

$$\forall n \in \mathbb{N} \quad 1 \cdot n = n$$

$$\forall m, n \in \mathbb{N} \quad (mn = n \Rightarrow m = 1)$$

Lösung 5.2. Null. Es gilt $0 = -0$ und die Null ist die einzige Zahl mit dieser Eigenschaft.

Lösung 5.4. Die Ergebnisse sind $-5, 6, 130, 46, 30, -6, 1, 40, 10, 80$ und 99 .

Lösung 5.5. Nur (iv), (viii) und (x) sind richtig. Ausführliche Lösungen finden Sie in diesem Video.

Lösung 5.6. Die Lösungen sind $1/3, 17/8, 29/30, 3/4$ und $1/42$.

Lösung 5.7. Die Lösungen sind $5a/18, (a + 2b)/10$ und $(a - b)/6$.

Lösung 5.8. Die Lösungen sind $(m + 2)/d, (a + b)/(1 - c), z/6$ und $a/3$.

Lösung 5.9. Mit *konkret* ist gemeint, dass Sie Zahlen einsetzen. Dafür reicht in diesem Fall schon ein so simples Beispiel wie $a = b = c = d = 1$. Dann steht links vom Gleichheitszeichen 2 und rechts 1.

Lösung 5.10. Keine der Gleichungen ist korrekt! In der Regel findet man ein Gegenbeispiel sofort, indem man für die Variablen irgendwelche natürlichen Zahlen einsetzt. Man muss schon sehr viel Pech haben, um zufällig Zahlen auszuwählen, für die die Gleichung stimmt.

Lösung 5.11. Der Kehrwert von -1 ist die ganze Zahl -1 .

Lösung 5.12. $A = \{3/7\}$.

Lösung 5.13. Nach dem besagten Korollar kann man die Menge als

$$B = \{x \in \mathbb{Q} : 3x - 4 = 0 \vee 2x + 5 = 0\}$$

schreiben. Es folgt (siehe Aufgabe 4.20)

$$B = \{x \in \mathbb{Q} : 3x - 4 = 0\} \cup \{x \in \mathbb{Q} : 2x + 5 = 0\} = \{4/3\} \cup \{-5/2\} = \{4/3, -5/2\}.$$

Lösung 5.14. Die Antwort ist $C = \{5\}$. Im Prinzip geht das so wie in Aufgabe 5.13, aber es ist zu beachten, dass C nur natürliche Zahlen enthalten soll. Von den beiden potentiellen Lösungen -5 und 5 kommt daher nur eine infrage.

Lösung 5.15. Da n eine natürliche Zahl sein soll, müssen die Faktoren $n + 3$ und $n + 4$ ebenfalls natürliche Zahlen sein. Es gibt aber nur einige wenige Möglichkeiten, 42 als Produkt zweier natürlicher Zahlen darzustellen, z. B. $1 \cdot 42$ oder $3 \cdot 14$. Und nur bei einem dieser Produkte ist die Differenz der beiden Faktoren 1, nämlich bei $6 \cdot 7$. Also folgt $n + 3 = 6$, $n + 4 = 7$ und damit $D = \{3\}$.

$n = -10$ würde die Gleichung auch lösen, aber -10 ist keine natürliche Zahl.

Lösung 5.16. Die Definition von f ist fehlerhaft, weil sich die Fälle $x \geq 3$ und $x \leq 3$ nicht gegenseitig ausschließen. Es ist nicht klar, ob $f(3) = 9$ oder $f(3) = 7$ gelten soll. In der Definition von g liegt dieselbe Überschneidung vor, aber „zufällig“ gilt $3 \cdot 2 = 2 \cdot 2 + 2$. Trotzdem sollte man so etwas vermeiden. Man hätte für den zweiten Fall einfach $x < 2$ schreiben können. Ob die Definition von h falsch ist, hängt davon ab, was erreicht werden sollte. Auf jeden Fall ist $h(x)$ für Zahlen zwischen 2 und 3 wie z. B. $x = 13/5$ nicht definiert.

Lösung 5.17. Die beiden Seiten sind dann nicht gleich, wenn x und y unterschiedliche Vorzeichen haben, z. B. $x = 1$ und $y = -1$.

Lösung 5.18. Ein einfaches Beispiel sind $a = 2$ und $b = 3$.

Lösung 5.19. Ist n positiv, dann ist 10^n im Dezimalsystem eine Eins gefolgt von n Nullen, z. B. $10^3 = 1000$. Für negative Exponenten erhält man logischerweise die Kehrwerte, also ist beispielsweise 10^{-2} ein Hundertstel.

Lösung 5.21. Die Lösungen sind 54, 100, 125, 216, $1/64$, 72, $1/81$, 108, $1/49$ und 1.

Lösung 5.22. Alle sind falsch. Richtig ist

$$(a^4)^7 = a^{4 \cdot 7} = a^{28},$$

$$a^8 - a = a \cdot a^7 - a \cdot 1 = a(a^7 - 1),$$

$$(-a)^2 = (-a) \cdot (-a) = a \cdot a = a^2 \quad \text{und}$$

$$a^{-5} = 1/a^5.$$

Lösung 6.1. Das einfachste Beispiel wären $A = \emptyset$ und $B = \{\emptyset\}$. Man könnte aber auch in Anlehnung an (6.1) $A = \{1, 2\}$ und $B = \{1, 2, A\}$ nehmen.

Lösung 6.3. Zunächst gilt $A \setminus \emptyset = A$ und $B \setminus \emptyset = B$. Das sollten Sie sich bereits in Aufgabe 2.4 klargemacht haben. $A \setminus \{\emptyset\}$ entsteht, indem man aus A das Element $\emptyset$ entfernt. Das Ergebnis ist $\{\{\emptyset\}\}$. Beachten Sie die doppelten Klammern! Die resultierende Menge hat ein Element und dieses Element ist selbst wiederum eine Menge. Bei $B \setminus \{\emptyset\}$ soll ebenfalls das Element $\emptyset$ entfernt werden. Da $\emptyset$ aber gar kein Element von B ist, ändert sich nichts: $B \setminus \{\emptyset\} = B$. Ebenso würde $B \setminus \{42\} = B$ gelten.

Diese Aufgabe wurde hauptsächlich deshalb aufgenommen, weil erfahrungsgemäß viele Studierende glauben, dass $\emptyset \in X$ für jede Menge X gilt. Das ist falsch! Es gilt $\emptyset \subseteq X$ für jede Menge X (siehe Aufgabe 2.4), aber das ist eine ganz andere Aussage. Wäre die leere Menge ein Element jeder Menge, dann wäre sie auch ein Element von $A \setminus \{\emptyset\}$, obwohl sie gerade entfernt wurde. Ist es nicht logisch, dass das nicht sein kann?

Lösung 6.4. Aussage (vi) ist wahr, siehe dazu Aufgabe 2.4. Alle anderen Aussagen sind falsch. Bei (ii) enthält die Menge links das Element $\{1\}$, die rechte aber nicht. Bei (i), (iii), (iv) und (v) enthält die linke Menge jeweils das Element 1 und die rechte nicht. Bei (vii) enthält die rechte Menge das Element $\emptyset$ und die linke nicht. Bei (viii) enthält umgekehrt die linke Menge das Element $\emptyset$ und die rechte nicht. Und bei (ix) enthält die linke Menge z. B. das Element $\{0, 1\}$ und die rechte nicht.

Lösung 6.5. Um diese Aufgabe richtig zu verstehen, muss man sich **daran erinnern**, dass $m, n \in A$ eine Abkürzung für $m \in A \wedge n \in A$ ist. Die Variablen m und n durchlaufen unabhängig voneinander die Menge A . Weil A drei Elemente hat, sind damit theoretisch $3 \cdot 3 = 9$ Produkte möglich. Allerdings ist die Multiplikation **kommutativ** und in Mengen gibt es keine Dopplungen. Bei $m = 3$ und $n = 2$ kommt beispielsweise dasselbe heraus wie bei $m = 2$ und $n = 3$. Daher ergibt sich $B = \{1, 2, 3, 4, 6, 9\}$. Die Antwort ist somit $|B| = 6$.

Ersetzt man 3 durch größere Zahlen, so kommt erschwerend hinzu, dass sich manche Produkte mit unterschiedlichen Faktoren bilden lassen: Bei $2 \cdot 6$ kommt etwa dasselbe heraus wie bei $3 \cdot 4$. Siehe dazu auch die „Computerlösung“ auf Seite 53.

Lösung 6.6. Sind sie auf diese Fangfrage hereingefallen? Die Menge $\{B\}$ ist ein Singleton. Dabei ist völlig unerheblich, wie B aussieht und ob es sich überhaupt um eine Menge handelt. Die Antwort ist schlicht und einfach $2^1 = 2$. Auch $\mathcal{P}(\{\mathbb{N}\})$ hat nur zwei Elemente, nämlich $\emptyset$ und $\mathbb{N}$.

Die Aufgabe zeigt, dass man in der Mathematik immer sehr genau hinschauen muss, weil in der Regel jedes Symbol von Bedeutung ist. Wäre es um $\mathcal{P}(B)$ gegangen, dann wäre die Antwort $2^{18} = 262\,144$ gewesen, aber $\mathcal{P}(B)$ und $\mathcal{P}(\{B\})$ sind eben zwei ganz unterschiedliche Objekte.

Lösung 6.7. Es stimmt *nie!* Seit Lemma 2.1 wissen wir, dass die leere Menge Teilmenge jeder Menge und deshalb Element jeder Potenzmenge ist. Darum enthält der Durchschnitt zweier Potenzmengen immer die leere Menge und ist deshalb nicht leer. (Wenn Sie diese Antwort verwirrt, dann sollten Sie vielleicht **den Abschnitt über Mengen** noch einmal studieren.)

Lösung 6.8. Für die erste Aussageform haben wir in Aufgabe 6.4 schon ein Gegenbeispiel gesehen: $A = \mathbb{N}$ und $B = \{0\}$. Für die zweite Aussageform kann man z. B. $A = \{0\}$ und $B = \{1\}$ nehmen. $\mathcal{P}(A \cup B)$ enthält dann die Menge $A \cup B = \{0, 1\}$, $\mathcal{P}(A) \cup \mathcal{P}(B)$ aber nicht.

Man beachte übrigens, dass $\mathcal{P}(A \cap B) = \mathcal{P}(A) \cap \mathcal{P}(B)$ immer gilt! Machen Sie sich das vielleicht anhand von Beispielen klar.

Lösung 6.9. In gewissem Sinne ist dies auch eine Fangfrage, bei der es um genaues Lesen geht. In der Aufgabe steht nicht, dass es um drei *verschiedene* Zahlen geht! Hat $\{a, b, c\}$ wirklich die Mächtigkeit 3, dann hat X die Mächtigkeit 2. Sind aber alle drei Zahlen identisch, dann hat X die Mächtigkeit 0. Ist auch die Antwort 1 möglich? Ja, wenn z. B. $a = b$ und $a \neq c$ gilt.

Lösung 6.10. Wenn Sie die Konzepte der Mengenlehre verstanden haben, sollte diese Aufgabe völlig unproblematisch sein.

$$\begin{array}{llllll} |A| = 0 & |B| = 1 & |C| = 1 & |D| = 3 & |E| = 1 & |F| = 2 \\ |G| = 2 & |H| = 4 & |I| = 3 & |J| = 3 & |K| = 1 & |L| = 1 \\ |M| = 1 & |N| = 3 & |O| = 2 & |P| = 2 & |Q| = 2 & |R| = 3 \end{array}$$

Lösung 6.11. Bei dieser Aufgabe geht es natürlich unter anderem auch wieder darum, die Mengenschreibweise zu verstehen. Wir haben:

$$\begin{array}{ll} A = \emptyset & |A| = 0 \\ B = \{0, 1, 2, 3, 4, 5, 6\} & |B| = 7 \\ C = \{7\} & |C| = 1 \\ D = \{0, 1, 2, 3, 4, 5\} & |D| = 6 \\ E = \{1, 2, 3, 6, 7, 14, 21, 42\} & |E| = 8 \\ F = \{0, 1, 2, \dots, 20\} & |F| = 21 \end{array}$$

Nach Mächtigkeit sortiert ergibt das A, C, D, B, E, F .

Lösung 6.12. Wenn Sie diese Aufgabe nicht lösen konnten, dann ist das nicht schlimm. Sie sollten es aber zumindest ernsthaft versucht haben. Das gilt ohnehin für alle Aufgaben! Und Ihnen ist dabei hoffentlich klar geworden, wonach überhaupt gefragt wurde.

Teilmengen von A , die mindestens eine Primzahl enthalten, sind zum Beispiel $\{3\}$, $\{2, 5, 10\}$ oder A selbst. Die Menge *aller* Teilmengen von A ist $\mathcal{P}(A)$. Da A 20 Elemente hat, hat $\mathcal{P}(A)$ nach Satz 6.1 2^{20} Teilmengen, also mehr als eine Million. Die kann man schwerlich alle einzeln durchgehen und auf Primzahlen überprüfen. Was nun? Bis hier sollten Sie ohne Hilfe gekommen sein. Hoffentlich haben Sie einen Zettel vor sich, auf den Sie exemplarisch ein paar der gesuchten Teilmengen gekritzelt haben.

Häufig ist es einfacher, nach dem *Gegenteil* zu fragen, und das ist auch hier so. Eine Teilmenge von A , die *nicht* mindestens eine Primzahl enthält, enthält *keine* Primzahl. Sie ist damit eine Teilmenge von

$$B = \{n \in A : n \notin \mathbb{P}\} = \{0, 1, 4, 6, 8, 9, 10, 12, 14, 15, 16, 18\}.$$

Wir können die Aufgabe damit so umformulieren, dass wir die Mächtigkeit von $\mathcal{P}(A) \setminus \mathcal{P}(B)$ wissen wollen. Daher ist die Antwort

$$|\mathcal{P}(A)| - |\mathcal{P}(B)| = 2^{|A|} - 2^{|B|} = 2^{20} - 2^{12}.$$

Das sind immer noch mehr als eine Million Mengen und ich habe die Zahl absichtlich nicht hingeschrieben. Um die geht es nämlich nicht. Bei den meisten Aufgaben reicht es, wenn Sie sie so weit lösen können, dass Sie für den letzten Schritt mit dem **einfachen Taschenrechner des Departments** auskommen würden. Dass Sie $2^{20} - 2^{12}$ eintippen und das Ergebnis abschreiben könnten, glaube ich auch so.

Lösung 6.13. Das kann man wie in Aufgabe 6.12 machen, wie Ihnen hoffentlich aufgefallen ist. $A = \{3, 4, 5, \dots, 29, 30\}$ hat 28 Elemente. Die Menge der Elemente von A , die durch 3 teilbar sind, ist

$$\{3, 6, 9, 12, 15, 18, 21, 24, 27, 30\}$$

und hat die Mächtigkeit 10. Die Menge der Elemente von A , die *nicht* durch 3 teilbar sind, hat daher die Mächtigkeit 18. Die Antwort ist demnach $2^{28} - 2^{18}$.

Lösung 7.2. Wenn Sie mit dem verlinkten Modell z. B. $424242 \bmod 17$ berechnen wollen, dann tippen Sie erst die Zahl 424242 ein, dann die Taste $(a \ b/c)$, dann die Zahl 17 und anschließend das Gleichheitszeichen. Ihnen ist hoffentlich klar, wie das angezeigte Ergebnis zu interpretieren ist.

Bei größeren Zahlen wie 42 424 242 klappt das allerdings evtl. nicht mehr. Da bleibt Ihnen bei primitiven Rechnern nichts anderes übrig, als erst zu dividieren und dann das Produkt aus dem **Divisor** und dem ganzzahligen Anteil des Quotienten vom Dividenten abzuziehen. Im konkreten Fall also erst $42\,424\,242 \div 17$ und dann $42\,424\,242 - 2\,495\,543 \times 17$. (Der Taschenrechner beherrscht zum Glück die Regel „Punkt- vor Strichrechnung“.)

Lösung 7.3. Nein. Die Aufzählung 1, -1 , a und $-a$ der trivialen Teiler könnte einen zu dieser Aussage verleiten. Aber die Zahlen 1 und -1 bilden eine Ausnahme, weil in diesen beiden Fällen $\{1, -1\} = \{a, -a\}$ gilt, es also nur zwei (triviale) Teiler gibt.

Lösung 7.4. Hier geht es natürlich darum, sich die Definition genau anzuschauen. Kann man eine entsprechende Zahl k finden? Die Antwort auf die erste Frage ist *ja*, denn man kann $k = 0$ wählen. *Jede* Zahl ist also Teiler der Null. Die Antwort auf die zweite Frage ist jedoch *nein*, denn für jede Zahl k gilt $k \cdot 0 = 0$, da kommt also nie 7 heraus. Null teilt nur sich selbst und keine andere Zahl.

Lösung 7.5. Man nennt eine Zahl *gerade*, wenn sie durch 2 teilbar ist. Dazu gehören also Zahlen wie 2, 42 oder -18 . Man nennt eine Zahl *ungerade*, wenn sie *nicht* gerade ist. Die Menge der geraden Zahlen ist also $\{2n : n \in \mathbb{Z}\}$ und die Menge der ungeraden Zahlen ist $\{2n + 1 : n \in \mathbb{Z}\}$, weil bei Division durch 2 außer 0 und 1 kein anderer Rest auftreten kann.

Die Frage steht insbesondere deshalb an dieser Stelle, weil immer mal wieder gefragt wird, ob die Null eine gerade Zahl sei. Ja, ist sie.

Die Eigenschaft einer ganzen Zahl, gerade oder ungerade zu sein, nennt man übrigens ihre *Parität*.

Parität

Lösung 7.6. (ii) und (iii) sind falsch, der Rest ist korrekt.

Lösung 7.7. Wegen des **Dezimalsystems** steht 230 für $2 \cdot 10^2 + 3 \cdot 10$. Weil 10 und 10^2 offensichtlich durch 10 teilbar sind, gilt das nach Lemma 7.3 auch für diese Linearkombination der beiden Zahlen.

Weil 230 durch 10 und 10 durch 5 teilbar ist, ist 230 nach Lemma 7.2 durch 5 teilbar. 5 ist trivialerweise durch 5 teilbar. Die Summe $235 = 230 + 5$ ist dann nach Lemma 7.3 ebenfalls durch 5 teilbar.

Weil 230 durch 10 und 10 durch 2 teilbar ist, ist 230 nach Lemma 7.2 durch 2 teilbar. 6 ist ebenfalls durch 2 teilbar. Die Summe $236 = 230 + 6$ ist dann nach Lemma 7.3 ebenfalls durch 2 teilbar.

Weil 200 durch 100 und 100 durch 4 teilbar ist, ist 200 nach Lemma 7.2 durch 4 teilbar. 36 ist ebenfalls durch 4 teilbar. Die Summe $236 = 200 + 36$ ist dann nach Lemma 7.3 ebenfalls durch 4 teilbar.

Lösung 7.8. Da kommt auch in allen drei Fällen 5 heraus. Das Vorzeichen spielt keine Rolle, weil der größte gemeinsame Teiler auf jeden Fall mindestens 1 ist, wenn nicht beide Zahlen null sind.

Lösung 7.9. In beiden Fällen ist das Ergebnis 42. Allgemein folgt aus der Definition sofort $\text{ggT}(a, a) = \text{ggT}(0, a) = |a|$. (Erinnern Sie sich an Aufgabe 7.4.)

Lösung 7.11. Das Ergebnis ist 19. Sie können sich analog zu Aufgabe 7.10 auch selbst Aufgaben mit mehr als zwei Zahlen stellen und diese lösen.

Lösung 7.13. Nein. Wenn $d = xa + yb$ gilt, dann kann man zur rechten Seite „gefahrlos“ $ba - ab = 0$ addieren. Das ergibt $d = (x + b)a + (y - a)b$. Und das kann man auch mehrfach oder alternativ mit vertauschten Rollen ($ab - ba$) machen. Es gibt also unendlich viele Möglichkeiten.

Im konkreten Beispiel hat $25 = -5a + 9b$ funktioniert, aber

$$25 = (-5 + 225)a + (9 - 400)b = 220a - 391b$$

tut es z. B. auch. Rechnen Sie nach!

Lösung 7.14. Man muss dafür im Wesentlichen nur einen Schritt modifizieren:

- (i) Setze A und B auf die Werte a und b .
- (ii) Gilt $A \cdot B = 0$, dann ist das Ergebnis das **Maximum** von A und B und der Algorithmus ist beendet.
- (iii) Ist $A > B$, dann ersetze A durch $A \bmod B$, anderenfalls ersetze B durch $B \bmod A$.
- (iv) Fahre fort mit Schritt (ii).

Die entscheidende Änderung ist, dass statt der Differenz der Rest berechnet wird. Damit bei dieser Rechnung nicht versehentlich durch null dividiert wird, sieht der Test in Schritt (ii) nun etwas anders aus. ($A \cdot B$ ist null, wenn mindestens einer der beiden Faktoren null ist.)

Ob diese Version schneller als das Original ist, wenn man sie auf einem Computer implementiert, ist übrigens gar nicht so leicht zu beantworten. Erstens kann je nach Prozessor die Division (also $\bmod$) zeitaufwendiger als die Subtraktion sein und zweitens spart man nicht in jedem Fall Schleifendurchläufe. Probieren Sie es z. B. mal mit 1008 und 623 und Sie werden sehen, dass man mit der Subtraktionsvariante nicht mehr Arbeit hat als mit der neuen.

Lösung 7.15. Wenn n und $n + 1$ einen gemeinsamen positiven Teiler d haben, dann gibt es nach Definition der Teilbarkeit Zahlen k_1 und k_2 mit $n = k_1 d$ und $n + 1 = k_2 d$. Offensichtlich muss $k_2 > k_1$ gelten. Bildet man nun die Differenz, so ergibt sich

$$1 = (n + 1) - n = k_2 d - k_1 d = (k_2 - k_1) d.$$

Weil $k_2 - k_1$ und d natürliche Zahlen sind, deren Produkt 1 ist, müssen beide den Wert 1 haben. Andere positive Teiler außer $d = 1$ gibt es also nicht.

Lösung 7.16. Dass er gekürzt ist. Ein Bruch liegt genau dann in gekürzter Form vor, wenn Zähler und Nenner teilerfremd sind. Wären sie es nicht, so könnte man den gemeinsamen Teiler herauskürzen.

Lösung 7.17. Wenn k ein gemeinsamer Teiler von a und b ist, dann teilt er nach Lemma 7.3 auch $xa + yb$, also 1. Da 1 aber nur die Teiler 1 und -1 hat, kommen als gemeinsame Teiler von a und b nur 1 und -1 infrage.

Lösung 7.18. Wenn man die Gleichung etwas umstellt, dann sieht sie so aus:

$$x!y! - x! - y! = 2 \tag{L-1}$$

Wenn $x \leq y$ gilt, dann sind alle drei Terme auf der linken Seite von (L-1) durch $x!$ teilbar, also ist die ganze linke Seite durch $x!$ teilbar und damit muss auch 2 durch $x!$ teilbar sein. Das geht aber nur, wenn $x! = 1$ oder $x! = 2$ gilt. Setzt man diese beiden Werte in (L-1) ein, so erhält man $y! - 1 - y! = 2$ (also $-1 = 2$) für $x! = 1$ und $2y! - 2 - y! = 2$ (also $y! = 4$) für $x! = 2$. Beides ist nicht möglich.

Wenn hingegen *nicht* $x \leq y$ gilt, so gilt $y < x$ und man kann genau wie eben argumentieren, wenn man die Rollen von x und y vertauscht. Es gibt also keine ganzen Zahlen, die die Gleichung erfüllen.

Lösung 7.19. Eine ausführliche Lösung gibt es in dem Video [Die Superzahl](#).

Lösung 8.2. Wenn man geduldig weitermacht (oder einen Computer die Arbeit machen lässt), kommt man auf $1 + 2 \cdot 3 \cdot 5 \cdot 7 \cdot 11 \cdot 13 = 30031 = 59 \cdot 509$. Entscheidend für den Beweis war aber auch nicht, ob 30031 eine Primzahl ist. Entscheidend ist, dass die Primteiler 59 und 509 von 30031 Primzahlen sind, die nicht in der Menge $\{2, 3, 5, 7, 11\}$ vorkommen.

Lösung 8.3. Bereits die Menge der geraden Zahlen ist unendlich und bis auf 2 sind alle geraden Zahlen zusammengesetzte Zahlen.

Lösung 8.5. Sinnvollerweise probieren Sie der Reihe nach zuerst die kleinsten Primzahlen durch. Teilweise kann man das (siehe Seite 57) schon direkt sehen. Da 32 892 eine gerade Zahl ist, teilen Sie also durch 2 und erhalten 16 446. Dabei ist zu beachten, dass ein Primfaktor mehrfach vorkommen kann: Im konkreten Fall können wir erneut durch 2 dividieren und erhalten 8 223. Diese Zahl ist durch 3 teilbar und wir erhalten 2 741. Das ist eine Primzahl und die Antwort wäre $32\,892 = 2^2 \cdot 3 \cdot 2\,741$ gewesen.

Lösung 8.6. Eine Zahl p ist nach Definition prim, wenn Sie sich durch keine positive Zahl (außer 1) teilen lässt, die kleiner als p ist. Wir müssen aber nicht *alle* Zahlen durchprobieren, die zwischen 1 und p liegen. Erstens müssen wir wegen der **Transitivität der Teilbarkeitsrelation** nur durch Primzahlen teilen. Zweitens, und das ist der eigentliche „Shortcut“, müssen wir nicht durch Zahlen teilen, die größer als $\sqrt{p}$ sind! Warum? Wenn p keine Primzahl ist, dann gibt es Zahlen q_1 und q_2 , die beide kleiner als p sind und für die $q_1 q_2 = p$ gilt. Wären die beide größer als $\sqrt{p}$, dann wäre ihr Produkt größer als $\sqrt{p} \cdot \sqrt{p} = p$. Das kann nicht sein, also ist mindestens eine der beiden Zahlen höchstens so groß wie $\sqrt{p}$ und das gilt „erst recht“ für die Primteiler dieser Zahl.

Im konkreten Fall können Sie also Ihren Taschenrechner $\sqrt{2741} \approx 52.3546$ berechnen lassen und wissen, dass Sie nur durch die 15 Primzahlen von 2 bis 47 teilen müssen. Die nächste Primzahl, 53, ist bereits zu groß.

Lösung 8.8. Wenn man ein paar Fälle durchgerechnet und sich ggf. auch noch die zugehörigen Primfaktorzerlegungen angeschaut hat, sollte man darauf kommen, dass immer

$$ab = \text{ggT}(a, b) \cdot \text{kgV}(a, b)$$

gilt. Das kleinste gemeinsame Vielfache erhält man aus den Primfaktorzerlegungen von a und b , indem man quasi das Vorgehen vor Aufgabe 8.7 „umkehrt“: Man braucht von allen Paaren von Primzahlpotenzen jeweils die mit dem *höheren* Exponenten.

Fortsetzung des konkreten Beispiels: $\text{kgV}(1008, 1080)$ ist $15120 = 2^4 \cdot 3^3 \cdot 5 \cdot 7$. Und $72 \cdot 15120$ ist dasselbe wie $1008 \cdot 1080$.

Lösung 8.9. ZAPPA wird zu

$$2^{26} \cdot 3^1 \cdot 5^{16} \cdot 7^{16} \cdot 11^1 = 11\,230\,071\,898\,079\,569\,920\,000\,000\,000\,000\,000.$$

1464843750 ist $2^1 \cdot 3^1 \cdot 5^{12}$ und das steht für AAL.

Nicht jede Zahl ist ein Code. Schon $10 = 2^1 \cdot 5^1$ ist keiner, denn zwischen 2 und 5 fehlt die zweite Primzahl 3.

Lösung 9.4. Die Ergebnisse der Beispielaufgaben sind $[2]_7$ und $[5]_6$. Beachten Sie, dass Sie statt $[25]_7 + [12]_7 = [25 + 12]_7$ auch $[25]_7 + [12]_7 = [4]_7 + [5]_7 = [4 + 5]_7$ rechnen können. Ebenso kann man $[11]_6 \cdot [7]_6 = [11 \cdot 7]_6$ durch $[11]_6 \cdot [7]_6 = [5]_6 \cdot [1]_6$ ersetzen. Man sollte „strategisch“ immer mit möglichst kleinen Werten arbeiten.

Wenn Sie sich selbst weitere Aufgaben stellen, können Sie Ihre Ergebnisse überprüfen, indem Sie in **Wolfram Alpha** so etwas eingeben wie

$$(25+12) \bmod 7 \quad \text{oder} \quad (11 \cdot 7) \bmod 6.$$

Lösung 9.5. Zuerst werden zwei Restklassen addiert. Das erste Gleichheitszeichen ist einfach die Anwendung der Definition (9.1). Beim zweiten Gleichheitszeichen wird die Kommutativität der „normalen“ Addition in $\mathbb{Z}$ verwendet. Das dritte

Gleichheitszeichen ist erneut eine Anwendung von (9.1), so dass ganz rechts wieder eine Summe zweier Restklassen steht, allerdings mit vertauschter Reihenfolge.

Lösung 9.7. Die Ergebnisse sind 5, 6 und 0.

Lösung 9.8. Eine Zahl ist genau dann gerade bzw. ungerade, wenn sie bei Division durch zwei den Rest null bzw. eins hat – siehe Aufgabe 7.5. Es geht also schlicht und einfach um das Rechnen in $\mathbb{Z}/2\mathbb{Z}$, das einfacher nicht sein könnte:

+	0	1
0	0	1
1	1	0

·	0	1
0	0	0
1	0	1

Der Additionstabelle entnimmt man z. B., dass die Summe zweier Zahlen gleicher Parität gerade ist, weil in der **Hauptdiagonale** nur Nullen stehen. Und so weiter.

Lösung 9.9. Schreiben wir G und U für *gerade* und *ungerade*, so ergibt sich für die möglichen Wegwischvorgänge der beteiligten Zahlen die folgende Situation:

	vorher	nachher	Bilanz G	Bilanz U
G	G	G	-1	±0
U	U	G	+1	-2
G	U	U	-1	±0

Das ist so zu lesen, dass beispielsweise (erste Zeile) beim Wegwischen zweier gerader Zahlen sich als Differenz wieder eine gerade Zahl ergibt, wodurch die Anzahl der geraden Zahlen um eins abnimmt, während sich die Anzahl der ungeraden Zahlen nicht ändert.

Die Spalte mit der Anzahl der geraden Zahlen ist grau, weil sie nicht relevant ist. Entscheidend ist, dass sich bei *jedem* Wischvorgang die Anzahl der ungeraden Zahlen entweder nicht ändert oder um zwei verringert. Das hat zur Folge, dass sie die ganze Zeit ungerade bleibt, weil sie am Anfang ungerade war. (Am Anfang waren es 1013 ungerade Zahlen.) Im letzten Schritt ist nur noch eine Zahl da und die kann nicht gerade sein, denn dann wäre die Anzahl der ungeraden Zahlen (null) gerade.

Lösung 9.10. Für den Datentyp `int` wurden 32 Bit verwendet und deshalb wird modulo 2^{32} gerechnet. Rechnen Sie es nach!

Der **arithmetische Überlauf**, der hier beschrieben wird, tritt in PYTHON oder COMMON LISP übrigens nicht ein, weil dort mit **Langzahlarithmetik** gerechnet wird.

Lösung 9.11. In 13! kommen die Faktoren 2, 5, und 10 vor, also kann man 13! durch 100 teilen. Darum muss als Ergebnis eine Zahl herauskommen, die in der Dezimaldarstellung am Ende zwei Nullen hat.

Lösung 9.12. Offenbar gilt $2^2 \bmod 5 = 4 \bmod 5 = 4 \neq 3 = 128 \bmod 5 = 2^7 \bmod 5$.

Satz 9.1 sagt nur etwas über Addition und Multiplikation aus. Von Potenzen ist dort nicht die Rede. Das falsche Argument aus der Aufgabenstellung führt eine

Berechnung im Sinne von „ $[2]_5^{[7]_5}$ “ durch, aber wir haben nicht definiert, was „Restklasse hoch Restklasse“ bedeuten soll, und man könnte das auch gar nicht sinnvoll definieren.

Wenn wir in der modularen Arithmetik Potenzen verwenden, dann sind diese **als Abkürzungen** gedacht. $[4]_7^3$ steht also beispielsweise für $[4]_7 \cdot [4]_7 \cdot [4]_7$. Daher gelten auch die aus diesen Abkürzungen folgenden einfachen Rechenregeln wie etwa $[a]_m^k \cdot [a]_m^n = [a]_m^{k+n}$.

Lösung 9.13. Alle Zahlen der Form $10^k - 1$ für natürliche Zahlen k sind offenbar durch 9 (und damit auch durch 3) teilbar. Es ergibt sich

$$\begin{aligned} [7164]_9 &= [7 \cdot 10^3 + 1 \cdot 10^2 + 6 \cdot 10^1 + 4 \cdot 10^0]_9 \\ &= [7]_9 \cdot [10^3]_9 + [1]_9 \cdot [10^2]_9 + [6]_9 \cdot [10^1]_9 + [4]_9 \cdot [10^0]_9 \\ &= [7]_9 \cdot [1]_9 + [1]_9 \cdot [1]_9 + [6]_9 \cdot [1]_9 + [4]_9 \cdot [1]_9 \\ &= [7]_9 + [1]_9 + [6]_9 + [4]_9 = [7 + 1 + 6 + 4]_9 \end{aligned}$$

und damit haben 7164 und die Quersumme von 7164 denselben Rest bei Division durch 9. Ersetzt man 9 durch 3, sieht die Rechnung genauso aus.

Lösung 9.14. Wenn man sich die geraden und die ungeraden Potenzen von zehn separat anschaut, so ergibt sich modulo 11 ein einfaches Muster, weil die Gleichung $10^2 = 100 = 9 \cdot 11 + 1$ gilt:

$$\begin{aligned} [10^2]_{11} &= [1]_{11} \\ [10^{2n}]_{11} &= [(10^2)^n]_{11} = ([10^2]_{11})^n = ([1]_{11})^n = [1]_{11} \\ [10^{2n+1}]_{11} &= ([10^2]_{11})^n \cdot [10]_{11} = [1]_{11} \cdot [10]_{11} = [10]_{11} = [-1]_{11} \end{aligned}$$

Lösung 9.15. Nennen wir die Zahl, die wir untersuchen wollen, m . Durch das Abtrennen der letzten Ziffer zerlegen wir diese Zahl in der Form $m = 10a + b$ und berechnen dann $a - 2b$. Ist nun $a - 2b$ durch 7 teilbar, so findet man eine Zahl k mit $a - 2b = 7k$. Multiplikation mit 10 liefert $10a - 20b = 70k$ bzw. (nach Addition von $21b$ auf beiden Seiten) die Darstellung $m = 10a + b = 70k + 21b$. Rechts steht jetzt eine Linearkombination der Zahlen 70 und 21, die beide durch 7 teilbar sind. Also ist die linke Seite, m , nach Lemma 7.3 auch durch 7 teilbar.

Ist umgekehrt m durch 7 teilbar, so findet man ein k mit $7k = m = 10a + b$ bzw. $b = 7k - 10a$. Damit folgt $a - 2b = a - 2 \cdot (7k - 10a) = 21a - 14k$ und wir haben $a - 2b$ als Linearkombination zweier Vielfacher von 7 dargestellt.

Lösung 9.16. Weil a/b im Allgemeinen keine ganze Zahl sein wird. Was soll beispielsweise $[3/4]_5$ bedeuten?

Lösung 9.17. Die Eins hat als Kehrwert sich selbst. Das ist offenbar in jedem Restklassenring so. Im konkreten Beispiel hat auch noch die Fünf einen Kehrwert, und zwar auch sich selbst: $5 \cdot 5 = 1$.

Lösung 9.19. Das erste Ergebnis ist 142. Im zweiten Fall kommt 8 heraus. Bei so kleinen Zahlen wie 13 geht es evtl. schneller, wenn man einfach die möglichen

Fälle im Kopf durchprobiert. Es kommen ja nur die elf Zahlen von 2 bis 12 infrage. (Und weil 13 eine Primzahl ist, wissen wir, dass es mit einer dieser Zahlen klappen muss.)

Lösung 9.21. $4/5$ ist eine Abkürzung für $4 \cdot 1/5$ und damit nach Aufgabe 9.19 für $4 \cdot 8$, also für 6.

Lösung 9.22. Die Voraussetzung für die folgenden Punkte ist natürlich, dass nach Satz 9.3 in jedem Restklassenkörper $\mathbb{Z}_p$ jedes von null verschiedene Element a einen Kehrwert a^{-1} hat. Nun zu den „Merkwürdigkeiten“:

- Ist $ax = b$ eine Gleichung in $\mathbb{Z}_p$ mit $a \neq 0$, so kann man beide Seiten mit a^{-1} multiplizieren und erhält die Lösung $x = a^{-1}b$.
- Gilt $ab = ac$ in $\mathbb{Z}_p$ mit $a \neq 0$, so kann man erneut beide Seiten mit a^{-1} multiplizieren und erhält $b = c$. Es ist also nicht möglich, dass beide Faktoren verschieden sind.

Man kann das auch so formulieren, dass die Produkte $1a, 2a, \dots, (p-1)a$ alle verschieden sind. Die Multiplikationstabelle bildet somit (abgesehen von der Zeile und der Spalte für die Null) ein lateinisches Quadrat.

- Gilt schließlich $ab = 0$ mit $a \neq 0$ so führt (zum dritten Mal) Multiplikation mit a^{-1} zur Gleichung $b = 0$. Mit anderen Worten: Ist a nicht null, dann muss b null sein.

Lösung 9.23. In $\mathbb{Z}_4$ ist -1 der Wert 3 und es gilt $3 \cdot 3 = 1$. In $\mathbb{Z}_7$ ist -1 der Wert 6 und es gilt $6 \cdot 6 = 1$. Und in $\mathbb{Z}_{13}$ ist das Ergebnis ebenfalls 1.

In jedem Restklassenring $\mathbb{Z}/n\mathbb{Z}$ gilt $(-1)^2 = 1$. Das kann man sich leicht überlegen, indem man die Berechnung in die ganzen Zahlen überträgt. Man berechnet dort $(n-1)^2 = n^2 - 2n + 1 = (n-2)n + 1$ und das ist offenbar eine Zahl, die den Rest 1 hat, wenn man durch n dividiert. (Eine alternative Begründung, die sogar noch allgemeiner ist, finden Sie in der Lösung zu Aufgabe 9.29.)

Damit folgt offensichtlich, dass (wie in den ganzen Zahlen) in jedem Restklassenring $(-1)^k = 1$ für gerade k und $(-1)^k = -1$ für ungerade k gilt.

Lösung 9.24. Ein PYTHON-Programm, das alle möglichen Fälle durchprobiert, könnte so aussehen:

```
x, a = 10000003, 0
while 2*a*a <= x:
    b = 0
    while a*a+b*b < x: b += 1
    if a*a+b*b == x:
        print(a,b)
    a += 1
```

Weil man keine Ausgabe erhält, erkennt man, dass es solche Zahlen a und b nicht gibt. Für wesentlich größere Zahlen als 10 000 003 würde das Programm allerdings ziemlich lange brauchen.

Die Begründung, warum es nicht geht, sieht so aus: Gäbe es Zahlen a und b mit der gewünschten Eigenschaft, dann müsste (9.4) auch für die Reste, also auch in $\mathbb{Z}/4\mathbb{Z}$ gelten. Weil 10 000 003 offensichtlich den Rest 3 hat, wenn man durch 4 teilt, müsste es also c und d in $\mathbb{Z}/4\mathbb{Z}$ mit $c^2 + d^2 = 3$ geben. Die Quadrate der Elemente 0, 1, 2 und 3 von $\mathbb{Z}/4\mathbb{Z}$ sind jedoch 0, 1, 0 und 1 und die Summe zweier solcher Quadrate kann nur $0 + 0 = 0$, $0 + 1 = 1$ oder $1 + 1 = 2$, aber niemals 3 sein.

Eine berühmte Verallgemeinerung dieser Aufgabe ist der [Zwei-Quadrate-Satz](#). Mehr dazu in meinem Video [Was ist Mathematik?](#) oder in meinem Buch [Pi und die Primzahlen](#) (das man über die [HAW-Bibliothek](#) kostenlos herunterladen kann).

Lösung 9.25. Der Fall $n = 1$ ist trivial; Y gewinnt immer, weil jede Zahl durch 1 teilbar ist. Ist hingegen n gerade oder die Zahl 5, dann kann X immer gewinnen. Dafür muss im ersten Zug lediglich eine der Ziffern 1, 3, 7 oder 9 in das Feld F eingetragen werden. (Erinnern Sie sich an die Teilbarkeitsregeln von Seite 57.) Für die verbleibenden Zahlen $n = 3$ oder $n = 9$ hat Y eine Gewinnstrategie. Mit dem letzten Zug kann nämlich offenbar immer erreicht werden, dass die [Quersumme](#) von m durch n teilbar ist.

Lösung 9.26. Es ist hilfreich, sich eine Tabelle der Reste anzuschauen.

		1	2	3	4	5	6	7	8	9
F	1	1	2	3	4	5	6	0	1	2
E	10	3	6	2	5	1	4	0	3	6
D	100	2	4	6	1	3	5	0	2	4
C	1000	6	5	4	3	2	1	0	6	5
B	10000	4	1	5	2	6	3	0	4	1
A	100000	5	3	1	6	4	2	0	5	3

Hier sind alle möglichen Spielzüge aufgeführt. Die hervorgehobene Zahl **2** bedeutet z. B., dass das Eintragen der Ziffer 5 im Feld C dem Rest $(5 \cdot 1000) \bmod 7 = 2$ entspricht. Jede eingetragene Ziffer „verschiebt“ also nach Satz 9.1 den sich am Ende ergebenden Gesamtrest $m \bmod 7$ um den entsprechenden Wert.

Auch ohne diese Tabelle hätte man nach Aufgabe 9.22 wissen können, dass in jeder Zeile in den ersten sieben Spalten alle Zahlen von 0 bis 6 erscheinen müssen. Das führt zu einer Gewinnstrategie für Y: Wählt X in seinem Zug eine Ziffer, die dem Rest r entspricht, so wählt Y im folgenden Zug eine Ziffer, die diesen Zug gewissermaßen neutralisiert – nämlich eine, die dem Rest $7 - r$ entspricht. Setzt X beispielsweise im ersten Zug eine Acht auf Feld D (das entspricht dem Rest 2), dann kann Y im nächsten Zug eine Zwei auf Feld C setzen. (Oder eine Eins auf Feld A. Für jedes Feld gibt es mindestens eine Möglichkeit, den Rest 5 zu erhalten.) Durch diese Strategie ist nach jedem Zug von Y der Gesamtrest null.

Lösung 9.27. Zunächst einmal kann man $a = 1$ für jedes n als Beispiel nehmen. Aber so ein Beispiel würde man sicher als trivial bezeichnen. Durch Probieren erhält man interessantere Beispiele wie $a = 8$ und $n = 9$.

Diese Aufgabe ist nebenbei ein schönes Beispiel für die wichtige Tatsache, dass $A \Rightarrow B$ und $B \Rightarrow A$ nicht äquivalent sind. Siehe Aufgabe 1.10 und die Bemerkungen am Ende des [Abschnitts über Aussagenlogik](#).

Lösung 9.28. Die Gleichung $a^{p-1} = 1$ kann man auch als $a \cdot a^{p-2} = 1$ hinschreiben und erkennt dadurch, dass a^{p-2} der Kehrwert von a ist.

Daraus folgt übrigens auch, dass die Umkehrung des kleinen Satzes von Fermat gilt: Wenn in einem Restklassenring $\mathbb{Z}/n\mathbb{Z}$ für *alle* Elemente a außer 0 die Gleichung $a^{n-1} = 1$ gilt, dann handelt es sich um einen Körper (und n ist eine Primzahl). Man vergleiche diese Aussage mit Aufgabe 9.27.

Lösung 9.29. Modulo 13 erhalten wir $4^2 = 3$ und $4^3 = 3 \cdot 4 = -1$. Damit folgt

$$4^{100} = 4^{99} \cdot 4 = (4^3)^{33} \cdot 4 = (-1)^{33} \cdot 4 = -1 \cdot 4 = 12 \cdot 4 = 9.$$

Den letzten Schritt kann man noch etwas abkürzen, wenn man sich klarmacht, dass in jedem Restklassenring $\mathbb{Z}/n\mathbb{Z}$ für jedes Element a der Zusammenhang $(-1) \cdot a = -a$ gilt. Das folgt aus

$$(-1) \cdot a + a = (-1) \cdot a + 1 \cdot a = (-1 + 1) \cdot a = 0 \cdot a = 0,$$

womit gezeigt ist, dass $(-1) \cdot a$ additiv invers zu a ist und damit quasi den „Namen“ $-a$ zu Recht trägt.

Dadurch hat man oben sofort $-1 \cdot 4 = -4 = 9$.

Lösung 9.30. Wir zerlegen zunächst:

$$\begin{aligned} 5^{37} &= 5 \cdot 5^{36} = 5 \cdot (5^{18})^2 = 5 \cdot ((5^9)^2)^2 = 5 \cdot ((5 \cdot 5^8)^2)^2 = 5 \cdot \left((5 \cdot (5^4)^2)^2 \right)^2 \\ &= 5 \cdot \left(\left(5 \cdot ((5^2)^2)^2 \right)^2 \right)^2 \end{aligned}$$

Und nun „von innen nach außen“:

$$\begin{aligned} 5^2 &= 7 \\ 5^4 &= (5^2)^2 = 7^2 = 4 \\ 5^8 &= (5^4)^2 = 4^2 = 7 \\ 5^9 &= 5 \cdot 5^8 = 5 \cdot 7 = 8 \\ 5^{18} &= (5^9)^2 = 8^2 = 1 \\ 5^{36} &= (5^{18})^2 = 1^2 = 1 \\ 5^{37} &= 5 \cdot 5^{36} = 5 \cdot 1 = 5 \end{aligned}$$

Lösung 9.31. Wie nehmen das Beispiel aus Aufgabe 9.30. Binär hat der Exponent 37 die Darstellung 100101. Den Rechenvorgang kann man mithilfe der folgenden Tabelle schematisch darstellen:

5^0	5^0	5^0	5^0	5^0	5^0
5^1	5^2	5^4	5^8	5^{16}	5^{32}
			5^9	5^{18}	5^{36}
					5^{37}
1	0	0	1	0	1

Man fängt links oben mit 5^0 an, also mit 1. Jeder Schritt nach rechts entspricht dem Quadrieren der linken Zahl; jeder Schritt nach unten entspricht dem Multiplizieren der oberen Zahl mit 5. (Die grauen Werte werden nicht benötigt und dienen nur der Illustration.)

Für jede Eins in der Binärdarstellung des Exponenten gibt es einen Schritt nach unten, für jede Ziffer der Binärdarstellung (bis auf die letzte) einen Schritt nach rechts. Die meiste „Arbeit“ machen also Zahlen wie 7, 31 oder 1023, deren Binärdarstellung aus lauter Einsen besteht. Und die Binärdarstellung liefert damit auch die gesuchte Obergrenze: Bei der Berechnung von a^b muss man maximal $2\lceil\log_2 b\rceil - 1$ Multiplikationen durchführen.²³ (Gemeint sind hier die [Aufrundungsfunktion](#) und der [Logarithmus](#) zur Basis 2.)

Nehmen wir als Beispiel $b = 123\,456\,789\,123\,456\,789\,123\,456\,789$. Würden Sie, um a^b zu berechnen, einfach in einer Schleife b -mal multiplizieren und würden Sie mithilfe eines [Supercomputers](#) eine Trillion (10^{18}) Multiplikationen pro Sekunde schaffen, dann würde Ihr Programm etwa 90 Jahre benötigen. Mit binärer Exponentiation würde derselbe Rechner hingegen nur Bruchteile einer Sekunde brauchen. (Im Vergleich zu den Zahlen, die in der Kryptographie verwendet werden, ist b klein.)

Lösung 10.1. Wir brauchen zunächst eine Lösung x_1 der ersten Kongruenz, die ein Vielfaches von $5 \cdot 7 = 35$ ist. Wir suchen also eine Zahl k_1 mit $[35k_1]_3 = [1]_3$. Da wir hier modulo 3 rechnen, können wir 35 durch 2 ersetzen und stattdessen $[2k_1]_3 = [1]_3$ lösen. Im allgemeinen Fall (z. B. in einem Programm) würden wir wie im Absatz hinter Satz 9.3 diese Gleichung mit dem erweiterten euklidischen Algorithmus lösen. Im konkreten Fall bietet sich jedoch simples Probieren an, weil für k_1 nur die Zahlen 0, 1 und 2 infrage kommen. $k_1 = 2$ tut es²⁴ und liefert uns $x_1 = 35 \cdot 2 = 70$. (In diesem Schritt rechnen wir nicht mehr modulo 3 und müssen mit 35 statt mit 2 multiplizieren.)

Für die beiden anderen Kongruenzen sind die jeweiligen Ansätze $[21k_2]_5 = [2]_5$ und $[15k_3]_7 = [4]_7$ bzw. vereinfacht $[k_2]_5 = [2]_5$ sowie $[k_3]_7 = [4]_7$. Damit erhalten

²³Hier wurde für a^1 eine Multiplikation gerechnet; nämlich $a^0 \cdot a = 1 \cdot a$. Bei den Angaben vor dieser Aufgabe wurde jeweils eine Multiplikation weniger gezählt. Für große Exponenten spielt dieser Unterschied natürlich keine Rolle.

²⁴Beachten Sie, dass der chinesische Restsatz eine Garantie dafür liefert, dass sich unter diesen drei Kandidaten eine Lösung befindet.

wir $k_2 = 2$, $k_3 = 4$, $x_2 = 21 \cdot 2 = 42$ und $x_3 = 15 \cdot 4 = 60$ und schließlich die Lösung $x = 70 + 42 + 60 = 172$.

Lösung 10.3. Anhand des vorgeführten Beispiels kann man sich überlegen, dass sich aus einer Lösung eine andere ergibt, wenn man das Produkt $m_1 \cdots m_n$ addiert oder subtrahiert. In Aufgabe 10.1 ist beispielsweise $172 - 3 \cdot 5 \cdot 7 = 67$ auch eine Lösung und 277 und -38 sind es auch. Es gibt also offensichtlich immer unendlich viele Lösungen und damit gibt es zum Beispiel auch eine eindeutig bestimmte kleinste Lösung, die nichtnegativ ist.

Lösung 10.4. Es ist einfach, ein unlösbares Kongruenzsystem anzugeben, wenn die Moduln identisch sind:

$$x \equiv 1 \pmod{3}$$

$$x \equiv 2 \pmod{3}$$

Aber das wäre geschummelt. Hier ist ein Beispiel mit verschiedenen Moduln:

$$x \equiv 1 \pmod{4}$$

$$x \equiv 2 \pmod{6}$$

Jedes x , das die erste Kongruenz löst, ist offenbar ungerade. Jedes x , das die zweite Kongruenz löst, ist gerade. Daher kann kein x beide Kongruenzen lösen.

Lösung 11.1. Addiert man beispielsweise die positive Zahl b zu a , so liegt die Summe $a + b$ weiter rechts als a . Der Abstand von $a + b$ zu a ist b . Addition einer negativen Zahl (bzw. Subtraktion einer positiven Zahl) bewirkt stattdessen eine Verschiebung nach links. Multiplikation mit einer positiven natürlichen Zahl kann man als wiederholte Addition im Sinne von $5a = a + a + a + a + a$ interpretieren. Die Division durch eine positive natürliche Zahl b entspricht einer Unterteilung in b gleich große Teile. (Der Quotient ist einer dieser Teile.) Und die Multiplikation mit dem Bruch a/b ist dann das Hintereinanderausführen der beiden gerade beschriebenen Operationen. Schließlich kann man die Multiplikation mit einer negativen Zahl c auffassen als Multiplikation mit der positiven Zahl $-c$ und anschließendem „Seitenwechsel“ des Produktes.

Wenn Sie beim Rechnen solche Vorstellungen im Kopf haben, werden Ihnen hoffentlich viele Dinge leichter fallen. Man kann sich auf diesem Wege auch die Rechenregeln aus Abschnitt 5 wie z. B. „minus mal minus ist plus“ (Lemma 5.4) klarmachen.

Lösung 11.3. Selbstverständlich müssen Ihre Gegenbeispiele nicht dieselben wie hier sein.

- (ii) und (iv) mit $a = b = -1$
- (vi) und (xiii) mit $a = 2$ und $b = -1$
- (x) mit $a = c = -1$ und $b = -2$
- (xi) mit $a = 2$ und $b = c = 1$
- (xii) mit $a = 1/2$ und $b = -1/2$

Lösung 11.4. Es werden die Aussagen, die *nicht* aus den Voraussetzungen folgen, jeweils mit einem Gegenbeispiel angegeben:

- (i) mit $c = 1/2$ und $d = 0$
- (iii) mit $a = 1$ und $d = 0$
- (v) mit $c = 1/2$ und $f = -1$
- (vi) mit $f = 0$
- (vii) mit $d = 1$ und $b = -1$
- (xi) mit $a = 1$, $b = 2$ und $d = 0$

Der Rest stimmt und lässt sich mit Lemma 11.3 begründen.

Erfahrungsgemäß haben in jedem Semester ein paar Studierende Schwierigkeiten mit Aussagen wie (viii). Deshalb sei zunächst daran erinnert, dass bei der Definition des Zeichens $\leq$ bereits darauf hingewiesen wurde, dass aus $ce < 0$ die Aussage $ce \leq 0$ folgt. Wenn Sie also meinen, dass (iv) korrekt ist, dann muss zwangsläufig auch (viii) richtig sein. Oder anders ausgedrückt: Es geht hier um die Aussage

$$\forall c \in \mathbb{Q} \forall e \in \mathbb{Q} (0 < c \wedge c < 1 \wedge e < 0) \Rightarrow (ce < 0 \vee ce = 0).$$

Wenn diese Aussage *nicht* wahr wäre, dann müssten Sie c und e angeben können, für die der Teil hinter den Quantoren falsch wird. Können Sie das?

Lösung 11.5. Es werden die Aussagen, die *nicht* wahr sind, jeweils mit einem Gegenbeispiel angegeben:

- (ii): Aus $a = b$ folgt immer $ac = bc$, aber die Umkehrung gilt nicht, z. B. mit $a = 1$, $b = 2$ und $c = 0$.
- (iii): Aus $a = b$ folgt immer $a^2 = b^2$, aber die Umkehrung gilt nicht, z. B. mit $a = 1$ und $b = -1$.
- (vi): Sei $c = -1$. Mit $a = 1$ und $b = 2$ gilt $a \leq b$, aber nicht $ac \leq bc$. Und mit $a = 2$ und $b = 1$ gilt $ac \leq bc$, aber nicht $a \leq b$.
- (vii): Aus $a > 0 \wedge b > 0$ folgt immer $a + b > 0$, aber die Umkehrung gilt nicht, z. B. mit $a = 2$ und $b = -1$.
- (viii): Aus $a > 0 \wedge b > 0$ folgt immer $ab > 0$, aber die Umkehrung gilt nicht, z. B. mit $a = b = -1$.

Die restlichen Äquivalenzen sind für alle rationalen Zahlen wahr.

Lösung 11.6. Zum Beispiel $a = 9$, $b = 1$, $c = 10$ und $d = 3$.

Lösung 11.8. A_1 ist weder nach unten noch nach oben beschränkt, weil $\mathbb{Z}$ eine Teilmenge von A_1 ist: Ist $k \in \mathbb{N}^+$, dann kann man k als $\frac{-k}{1}$ und $-k$ als $\frac{-k}{1}$ darstellen, wobei in beiden Fällen die Summe von Zähler und Nenner auf jeden Fall unterhalb von einer Million liegt.

A_2 ist als Teilmenge von $\mathbb{N}$ nach unten beschränkt, aber nicht nach oben, weil z. B. alle geraden Zahlen, die größer als 2 sind, zu A_2 gehören.

A_3 ist beschränkt. Nach unten ist die Menge durch 0 beschränkt, nach oben durch 1 000 000.

A_4 ist beschränkt. Die Reste bei Division durch b liegen nach Definition zwischen 0 und b und damit ist A_4 die Menge $[41] \cup \{0\}$.

A_5 ist beschränkt. Es handelt sich um die Menge $\{x \in \mathbb{Q} : 0 < x \leq 1\}$. (Siehe Korollar 11.4.)

A_6 ist beschränkt. Man kann die Schranken 0 (unten) und 42 (oben) direkt ablesen. Dass A_6 keine natürlichen Zahlen enthält, spielt keine Rolle.

A_7 ist als Menge von natürlichen Zahlen nach unten beschränkt, aber nicht nach oben, weil alle natürlichen Zahlen, die größer als 42 sind, zu A_7 gehören.

A_8 ist beschränkt. Es handelt sich um die Menge A_4 .

Nach dem kleinen Satz von Fermat ist A_9 die Menge $\{0, 1\}$, also beschränkt.

A_{10} ist beschränkt. Es handelt sich um die Menge A_5 .

A_{11} ist nach oben durch 1 beschränkt, nach unten jedoch nicht, denn $x^3 \leq x$ gilt z. B. für jede negative ganze Zahl x .

A_{12} ist beschränkt. Es handelt sich um die Menge $\{q \in \mathbb{Q} : 0 \leq q < 1\}$.

Lösung 11.9. Man kann die größere der beiden Zahlen $|a|$ und $|b|$ nehmen. Sind beide gleich groß, ist es natürlich egal, welche man nimmt. (Beachten Sie, dass Sie nicht wissen, ob a bzw. b positiv oder negativ sind!)

Lösung 11.10. Man könnte gleich das erste Beispiel $B_1 = \{r \in \mathbb{Q} : r < 3\}$ hinter der Definition des Begriffs *beschränkt* nehmen.

Falls Ihnen so ein Beispiel nicht eingefallen ist, dann machen Sie sich noch einmal die hinter Lemma 11.5 angesprochene unintuitive Eigenschaft klar. Alle Elemente von B_1 sind kleiner als 3, aber zu jedem Element b von B_1 gibt es unendlich viele andere Elemente von B_1 , die größer als b sind.

Lösung 11.11. Da gibt es natürlich viele. Man könnte der Einfachheit halber für beide Fragen die Menge $\mathbb{Z}$ nehmen.

Lösung 11.12. Auch hier gibt es viele mögliche Antworten. Zum Beispiel die Mengen $\mathbb{N}$ und $\mathbb{P}$.

Lösung 11.13. Wir haben den Algorithmus in der Form einer Tabelle mit zwei Spalten aufgeschrieben. In jeder Zeile ist eine der beiden Zahlen kleiner als die über ihr. Daher wird die Zeilensumme in jedem Schritt kleiner. Wenn der Algorithmus nie terminieren würde, dann gäbe es unendlich viele Zeilen und die Menge der Zeilensummen wäre eine Menge natürlicher Zahlen ohne Minimum. Wegen der Wohlordnung der natürlichen Zahlen kann es so eine Menge nicht geben.

Mit anderen Worten: Wann immer es in einer Schleife in einem Programm eine ganzzahlige Variable gibt, deren Wert erstens in jedem Schritt kleiner und zweitens nie negativ wird, können Sie sich sicher sein, dass es sich nicht um eine Endlosschleife handelt.

Lösung 12.1. $x^2 = 13$ hat keine rationale Lösung, denn 13 ist keine Quadratzahl. $x^2 = 121$ hat die rationale Lösung 11. $x^2 = 45/80$ hat die rationale Lösung $3/4$, denn $45/80$ hat gekürzt die Form $9/16 = 3^2/4^2$. $x^5 = 1/32$ hat die rationale Lösung $1/2$, weil $32 = 2^5$ (und trivialerweise $1 = 1^5$) gilt. $x^3 = 27$ hat die rationale Lösung 3. $x^3 = 29$ hat keine rationale Lösung: So eine Lösung müsste eine natürliche Zahl sein, aber 3^3 ist kleiner als 29 und 4^3 größer. $x^4 = 16/125$ hat keine rationale Lösung: Zwar gilt $16 = 2^4$, aber 125 ist nicht die vierte Potenz einer natürlichen Zahl. $x^3 = 3/81$ hat die rationale Lösung $1/3$, denn $3/81$ ist gekürzt der Bruch $1/27 = 1/3^3$.

Lösung 12.2. Für eine rationale Lösung p/q mit natürlichen Zahlen p und q müsste $p^3/q^3 = 29/1$ gelten. Wenn p und q teilerfremd sind, folgt $q^3 = 1$, also $q = 1$ und damit $p/q = p$.

Lösung 12.4. Die Ergebnisse sind $8/33$ und $10/27$.

Lösung 12.5. Es ergibt sich $9/9$, also 1. Wenn Sie das trotz der mathematischen Herleitung nicht so richtig „glauben“ können, dann schauen Sie sich das Video *Das Rätsel des Kontinuums* an.

Lösung 12.7. Das Ergebnis ist $42/55$. Siehe auch Aufgabe 12.8.

Lösung 12.9. Wäre die Zahl $r = 2 + \pi$ rational, so wäre $\pi = r + (-2)$ nach Satz 5.5 als Summe zweier rationaler Zahlen ebenfalls rational. Analog argumentiert man mit dem Produkt für $s = 2\pi$ mit $\pi = s \cdot 1/2$.

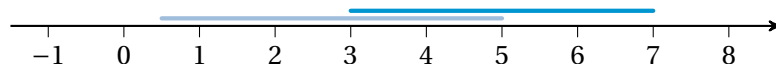
Lösung 12.10. Natürlich geht das. Die Summe von $-\pi$ und $2 + \pi$ ist beispielsweise 2, das Produkt von $1/\pi$ und 2π ist ebenfalls 2. Dass $-\pi$ und $1/\pi$ irrational sind, begründet man wie in Aufgabe 12.9.

Lösung 12.11. $[42, 42] = \{42\}$ und $[42, 42) = \emptyset$. Das sind also *unechte* Intervalle. Und offenbar ist jedes Singleton $\{r\}$ ein Intervall, nämlich $[r, r]$.

Lösung 12.12. $\mathbb{R}^+$ ist die Menge $(0, \infty)$ und $\mathbb{R}^-$ ist $(-\infty, 0)$.

Lösung 12.13. Hoffentlich haben Sie nicht gedacht, dass diese Menge fünf Elemente hat. In Lemma 12.3 steht ja auch explizit, dass alle echten Intervalle unendlich viele Elemente haben. Ich habe diese Aufgabe aufgenommen, weil es eine offenbar unausrottbare Fehlinterpretation gibt, nach der manche Studierende eine Menge wie $[1, 5]$ als die Menge der *ganzen* Zahlen zwischen 1 und 5 interpretieren. Das ist aber schlicht und einfach falsch und widerspricht der Definition, in der explizit steht, dass $[1, 5]$ die Menge aller *reellen* Zahlen zwischen 1 und 5 ist!

Lösung 12.14. Grundsätzlich kann es für solche Fragen hilfreich sein, sich die beteiligten Intervalle als Abschnitte der Zahlengeraden zu skizzieren, z. B. so:



Der Durchschnitt ist dann sofort erkennbar. Lediglich bei den Intervallgrenzen muss man explizit prüfen, ob sie in der Schnittmenge liegen oder nicht. Die Er-

gebnisse in diesem Fall sind in dieser Reihenfolge $(3, 5)$, $[3, 5]$, $(3, 7]$, $[3, 3] = \{3\}$ und $[3, 3) = \emptyset$.

Lösung 12.15. Die Vereinigung von $A = [1/3, 42)$ und $[42, 43)$ ist das Intervall $[1/3, 43)$. Die Vereinigung von A und $(42, 43)$ ist kein Intervall, weil die Zahl 42 „fehlt“ und damit für eine Lücke sorgt.

Lösung 12.16. Nein. $[0, 3] \setminus [1, 2]$ ist beispielsweise kein Intervall. Und diese Menge lässt sich auch nicht „kürzer“ schreiben, lediglich anders, etwa als $[0, 1) \cup (2, 3]$.

Lösung 12.17. $[40, 42] \setminus [40, 41)$ ist das Intervall $[41, 42]$. $[40, 41] \setminus \mathbb{N}$ ist das Intervall $(40, 41)$. $[40, 42] \setminus [40, 42)$ ist das Intervall $\{42\}$, das aber kein *echtes* Intervall ist. $[40, 41] \setminus \mathbb{Q}$ ist kein Intervall, denn diese Menge hat unendlich viele „Lücken“. (Zwischen 40 und 41 befinden sich unendlich viele rationale Zahlen, die alle entfernt wurden.) $[40, 41] \cap [41, 43]$ ist das Intervall $\{41\}$, das ebenfalls kein *echtes* Intervall ist. $(40, 42) \setminus \mathbb{N}$ ist die Menge $(40, 41) \cup (41, 42)$ mit der „Lücke“ 41.

Lösung 12.18. Die Antworten:

$$\sqrt{(-3)^2} = \sqrt{9} = 3$$

$$\sqrt[3]{27} = 3$$

$$\sqrt[3]{1/8} = 1/\sqrt[3]{8} = 1/2$$

$$\sqrt[4]{32}/\sqrt[4]{4} = \sqrt[4]{32}/\sqrt[4]{4} = \sqrt[4]{32/4} = \sqrt[4]{8} = \sqrt[4]{16/2} = \sqrt[4]{16}/\sqrt[4]{2} = 2/\sqrt[4]{2} = 2$$

$$\sqrt{4/9} = \sqrt{4}/\sqrt{9} = 2/3$$

$$\sqrt{2/3} \cdot \sqrt{6} = \sqrt{2/3 \cdot 6} = \sqrt{4} = 2$$

$$16^{1/4} = \sqrt[4]{16} = 2$$

$$(\sqrt[3]{5})^2 \cdot \sqrt[6]{25} = (\sqrt[3]{5})^2 \cdot \sqrt[3]{\sqrt{25}} = (\sqrt[3]{5})^2 \cdot \sqrt[3]{5} = (\sqrt[3]{5})^3 = 5$$

$$25^{-1/2} = 1/25^{1/2} = 1/\sqrt{25} = 1/5$$

Lösung 12.19. Die zweite Aussage ist korrekt, weil es sich um die Definition des Wurzelzeichens handelt: $\sqrt{a}$ ist eine Zahl, deren Quadrat a ist. Die erste Aussage ist falsch. Man könnte als Beispiel den Anfang von Aufgabe 12.18 mit $a = -3$ nehmen. Eine Gleichung, die für alle $a \in \mathbb{R}$ stimmt, ist $\sqrt{a^2} = |a|$.

Lösung 12.20. Sie sind hoffentlich auf diese Ergebnisse gekommen:

$$A = \{0\} \quad B = \mathbb{N} \quad C = \mathbb{Q} \quad D = \mathbb{Q}$$

$$E = \mathbb{Q} \quad F = \emptyset \quad G = \mathbb{N} \quad H = \mathbb{Z}$$

Lösung 12.21. Die Paare sind $A = C = \mathbb{R}$, $B = E = \mathbb{N}$, $D = I = \{0\}$, $F = J = \emptyset$ und $G = H$. Unter anderem muss man dafür wissen, dass $\sqrt{8} = 2\sqrt{2}$ nach dem **Satz des Theaitetos** keine rationale Zahl ist und dass 0 und 1 die einzigen Zahlen sind, die ihr eigenes Quadrat sind: $0 \cdot 0 = 0$ und $1 \cdot 1 = 1$.

Lösung 12.22. Die Paare sind $A = G$, $B = F = \mathbb{N}$, $C = D = \mathbb{N}^+$ und $E = H = \emptyset$. A ist die Menge der nichtnegativen rationalen Zahlen.

Lösung 12.23. Die Paare sind $A = (-2, 2) = H$, $B = \mathbb{R} = G$, $C = E$, $D = \emptyset = I$ und $F = J$.

Lösung 13.1. Die Antworten für π sind 3.142, 3.1416 und 3.1415927. Bei 3/80 kommt 0.037 oder 0.038 heraus. Sowohl bei kaufmännischem als auch bei mathematischem Runden würde man sich für 0.038 entscheiden müssen.

Lösung 13.2. Die Frage stellt sich offenbar nur, wenn vor einer Fünf gerundet wird, auf die nur noch Nullen folgen. Bei irrationalen Zahlen wie $\sqrt{2}$ kann dieser Fall nicht eintreten.

Lösung 13.3. Beispielsweise 0.0996. Gerundet ist das 0.100.

Lösung 13.5. Ein besonders einfaches Beispiel ist $x = y = 1$ mit $\tilde{x} = \tilde{y} = 1.1$. Die absoluten Fehler sind in beiden Fällen 0.1, die relativen Fehler liegen jeweils bei 10 %. Die „falsche“ Summe 2.2 weicht um 0.2 von der richtigen ab, das „falsche“ Produkt 1.21 ist 21 % größer als das richtige.

Lösung 14.1. Die Lösungen sind

$$\begin{aligned} z + w &= \frac{23}{21} - \frac{15}{8} \cdot i, \\ z - w &= \frac{5}{21} + \frac{25}{8} \cdot i, \\ zw &= \frac{207}{112} - \frac{235}{168} \cdot i \quad \text{und} \\ (\sqrt{2} - 3i) + (5 - \pi i) &= (5 + \sqrt{2}) - (3 + \pi)i. \end{aligned}$$

Bei der letzten Aufgabe ging es in erster Linie darum, dass es für Ausdrücke wie $5 + \sqrt{2}$ oder $3 + \pi$ keine einfachere Schreibweise gibt. Relevant ist, dass das Ergebnis so hingeschrieben wird, dass beim Ergebnis auf den ersten Blick Real- und Imaginärteil erkennbar sein sollen. $(5 + \sqrt{2}) + (-3 - \pi)i$ wäre übrigens auch richtig, ist aber etwas weniger elegant. Ihnen ist hoffentlich klar, dass das zweite Klammernpaar *nicht* weggelassen werden kann. Das vordere Paar ist hingegen rein formal überflüssig, hilft aber beim Lesen der Zahl.

Lösung 14.2. $\operatorname{Re}(z) = 3$ und $\operatorname{Im}(z) = -\sqrt{5}$. Eigentlich ganz einfach. Der Sinn dieser Aufgabe ist, auf zwei typische Fehler hinzuweisen: Erstens ist der Imaginärteil *nicht* $\sqrt{5}$, sondern $-\sqrt{5}$, weil eine komplexe Zahl *nach Definition* immer die Form $a + bi$ hat und b in diesem Fall die negative Zahl $-\sqrt{5}$ ist. Zweitens ist der Imaginärteil *nicht* $-\sqrt{5}i$. Der Imaginärteil ist immer nur die *reelle* Zahl b , die *vor* der imaginären Einheit steht.

Lösung 14.3. $z/w = -77/388 + 175i/582$ und $w/z = -(396 + 600i)/259$.

Lösung 14.5. Das ist nun wirklich einfach.

$$\begin{aligned} \overline{w + z} &= \overline{(a + c) + (b + d)i} = (a + c) - (b + d)i = a - bi + c - di \\ &= \overline{a + bi} + \overline{c + di} = \overline{w} + \overline{z} \\ \overline{w \cdot z} &= \overline{(ac - bd) + (ad + bc)i} = (ac - bd) - (ad + bc)i = (a - bi) \cdot (c - di) \\ &= \overline{a + bi} \cdot \overline{c + di} = \overline{w} \cdot \overline{z} \end{aligned}$$

In beiden Fällen kommt also jeweils dasselbe heraus.

Lösung 14.6. Ist $z = a + bi$, so erhält man

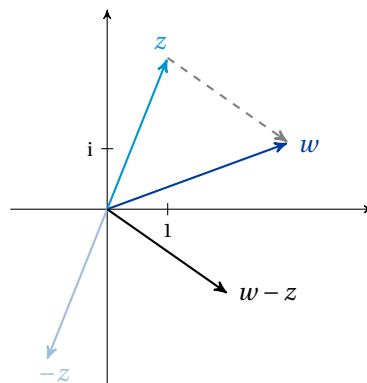
$$z + \bar{z} = (a + bi) + (a - bi) = 2a = 2\operatorname{Re}(z) \quad \text{und}$$

$$z - \bar{z} = (a + bi) - (a - bi) = 2bi = 2i\operatorname{Im}(z).$$

Lösung 14.7. Hat z die Koordinaten (a, b) , so hat die Zahl $-z$ die Koordinaten $(-a, -b)$. Den Ort von $-z$ erhält man also durch **Punktspiegelung** von z am Ursprung.

Lösung 14.8. Hat z die Koordinaten (a, b) , so hat die Zahl $\bar{z}$ die Koordinaten $(a, -b)$. Den Ort von $\bar{z}$ erhält man also durch **Achsen Spiegelung** von z an der reellen (horizontalen) Achse.

Lösung 14.9. Die Skizze geht mithilfe von Aufgabe 14.7 aus der für die Addition hervor, wenn man sich daran erinnert, dass $w - z$ für $w + (-z)$ steht.



Beachten Sie, dass wegen $w = z + (w - z)$ der (verschobene) $w - z$ -Pfeil von der Spitze von z auf die von w zeigt. (In der Skizze ist das der gestrichelte graue Pfeil.)

Lösung 14.10. In Gleichung (14.1). Das heißt also, dass allgemein $|z| = \sqrt{z\bar{z}}$ für jede komplexe Zahl z gilt.

Lösung 14.11. $|i| = \sqrt{0^2 + 1^2} = 1$.

Lösung 14.12. $|-z| = |-a - bi| = \sqrt{(-a)^2 + (-b)^2} = \sqrt{a^2 + b^2} = |z|$. Angesichts von Aufgabe 14.7 ist geometrisch auch klar, warum in beiden Fällen dasselbe herauskommen muss.

Lösung 14.13. Für reelle Zahlen kommt in beiden Fällen dasselbe heraus. (Wenn dem nicht so wäre, wäre die Definition auch sehr unglücklich.) Setzt man eine reelle Zahl a in die neue (komplexe) Definition ein, dann erhält man

$$|a| = |a + 0i| = \sqrt{a^2 + 0^2} = \sqrt{a^2}$$

und das ist der reelle Betrag.²⁵ Wie schon mehrfach erwähnt ist auch für reelle Zahlen der Betrag also der Abstand von der Null.

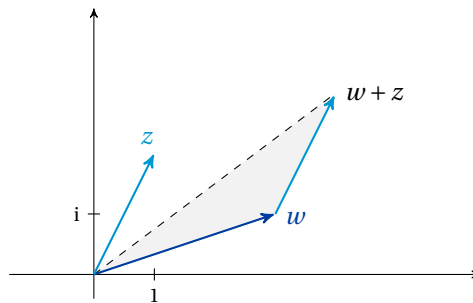
²⁵Achtung: Ein häufig anzutreffender Denkfehler ist, dass $\sqrt{a^2}$ dasselbe wie a ist. Dass das falsch ist, kann man sehen, wenn man für a eine negative Zahl einsetzt. Siehe dazu auch Aufgabe 12.19.

Lösung 14.14. Wenn man das Rechnen nicht verlernt hat, ist das Routine:

$$\begin{aligned}
 |wz| &= |(a+bi)(c+di)| = |(ac-bd) + (ad+bc)i| \\
 &= \sqrt{(ac-bd)^2 + (ad+bc)^2} \\
 &= \sqrt{a^2c^2 - 2acbd + b^2d^2 + a^2d^2 + 2adbc + b^2c^2} \\
 &= \sqrt{a^2c^2 + b^2d^2 + a^2d^2 + b^2c^2} \\
 |w| \cdot |z| &= |a+bi| \cdot |c+di| = \sqrt{a^2+b^2} \cdot \sqrt{c^2+d^2} \\
 &= \sqrt{(a^2+b^2) \cdot (c^2+d^2)} = \sqrt{a^2c^2 + b^2c^2 + a^2d^2 + b^2d^2}
 \end{aligned}$$

Aufgefallen sein sollte Ihnen, dass in beiden Fällen dasselbe herauskommt.

Lösung 14.15. Ihre Skizze sollte in etwa so aussehen:



In der Dreiecksungleichung geht es um das grau eingefärbte Dreieck und die geometrisch evidente Aussage, dass der direkte Weg (die gestrichelte Linie) von 0 zum Ziel $w+z$ kürzer ist als der „Umweg“ über w (entlang der blauen Pfeile).

Lösung 14.16. Weil der Realteil von z negativ und der Imaginärteil positiv ist, muss das Argument zwischen 90 und 180 Grad liegen. (Machen Sie sich eine Skizze, wenn Ihnen das nicht klar ist!) Das Argument von z^2 ist doppelt so groß und liegt daher zwischen 180 und 360 Grad (bzw. „offiziell“ zwischen -180 und 0 Grad, aber das ist für die geometrische Lage natürlich irrelevant). Deshalb liegt z^2 unterhalb der reellen Achse und kann *nicht* in der oberen Hälfte liegen.

Lösung 14.17. Das läuft analog zur letzten Aufgabe. Das Argument von z liegt im Bereich von 0 bis -90 Grad. (Skizze!) Das Argument von z^3 liegt daher zwischen 0 und -270 Grad und damit *nicht* im oberen rechten **Quadranten**.

Lösung 14.18. Am Anfang wurde durch p geteilt. Das ist natürlich nur möglich, wenn p nicht null ist. Ansonsten hätte auch die folgende Multiplikation mit p^2 keinen Sinn ergeben. Denn Fall kann man jedoch vorab klären: Wäre p null, dann wäre $w = qi$ rein imaginär und damit $w^2 = -q^2$ eine negative reelle Zahl.

Tatsächlich sollte man für das Lösen einer Gleichung der Form $w^2 = z$ folgendermaßen vorgehen:

- (i) Ist $z = 0$, so gibt es nur die eine Lösung $w = 0$.
- (ii) Ist z eine positiv reelle Zahl, so sind die beiden Lösungen $\sqrt{z}$ und $-\sqrt{z}$.

- (iii) Ist z eine negative reelle Zahl, so sind die beiden Lösungen $i\sqrt{-z}$ und $-i\sqrt{-z}$.
- (iv) Ist z eine echt komplexe Zahl, so funktioniert der eben vorgeführte Ansatz $w = p + qi$, wobei p und q nicht null sind und die zwischendurch auftretende quadratische Gleichung immer genau eine positive reelle Lösung hat.

Lösung 14.21. Es gibt vier Lösungen. Mit der Substitution $x^2 = w$ erhält man zunächst die Gleichung $w^2 = 16$, die die beiden Lösungen $w_1 = 4$ und $w_2 = -4$ hat. Die Rücksubstitution $x^2 = w_1 = 4$ liefert die beiden Lösungen $x_1 = 2$ und $x_2 = -2$. Die Rücksubstitution $x^2 = w_2 = -4$ liefert zwei weitere Lösungen, nämlich $x_3 = 2i$ und $x_4 = -2i$.

Dass es diese vier Lösungen gibt, kann man sich auch klarmachen, indem man sie in die gaußsche Zahlenebene einzeichnet und sich überlegt, was geometrisch passiert, wenn man ihre vierte Potenz bildet. Damit wird dann zumindest anschaulich klar, dass die Gleichung $x^4 = w$ für jede komplexe Zahl $w \neq 0$ vier Lösungen hat, die alle auf einem Kreis mit dem Radius $\sqrt[4]{|w|}$ liegen, dessen Mittelpunkt die Null ist.

Lösung 14.22. Das kann nicht sein. Aus $a + bi = 0$ folgt $a = -bi$. Mit $b = 0$ würde daher auch $a = 0$ folgen. Wäre jedoch b nicht null, so könnte man durch b dividieren und erhielte $-a/b = i$. Dann stünde jedoch links die reelle Zahl $-a/b$ und rechts die Zahl i , die wegen $i^2 = -1$ sicher nicht reell ist.

Lösung 15.1. Die Dopplungen verschwinden. Beispielsweise ist $2 \cdot 3$ dasselbe wie $3 \cdot 2$. Aber $(2, 3)$ ist eben *nicht* dasselbe wie $(3, 2)$. Das ist ja gerade der Grund, warum Paare definiert wurden.

Lösung 15.2. $\{0\} \times M$ ist $\{(0, 23), (0, 42), (0, \pi)\}$. $\emptyset \times M$ ist die leere Menge, denn in dieser Menge befinden sich nach Definition Paare, deren erste Komponente ein Element der leeren Menge ist. Da die leere Menge keine Elemente hat, kann es solche Paare nicht geben.

Etwas anderes hätte auch gar nicht herauskommen dürfen, denn nach Lemma 15.1 hat $\emptyset \times M$ die Mächtigkeit null.

Lösung 15.3. Die Antwort ist *nein* und als Begründung reicht bekanntlich nur ein Gegenbeispiel, etwa $A = \{1\}$ und $B = \{2\}$. (Wenn Ihnen das nicht sofort klar ist, dann schreiben Sie $A \times B$ und $B \times A$ auf.) Allgemein gilt $A \times B = B \times A$ dann und *nur* dann, wenn $A = B$ gilt oder eine der beiden Mengen leer ist.

Lösung 15.4. Den ersten Teil kann man natürlich *nicht* realisieren. Nach Lemma 15.1 müsste $|A|$ dann die Mächtigkeit $\sqrt{42}$ haben. Aber die Mächtigkeit einer (endlichen) Menge muss offensichtlich eine (nichtnegative) ganze Zahl sein und $\sqrt{42}$ ist *keine ganze Zahl*.

Für den zweiten Teil gibt es viele Möglichkeiten, z. B. $A = \{0, 10, 20, 30, 40, 50, 60\}$ und $B = \{0, 1, 2, 3, 4, 5\}$.

Lösung 15.5. Offenbar gilt $|A| = 6$ und damit $|B| = 2^6$ nach Satz 6.1. Nach Lemma 15.1 folgt $|C| = 2^6 \cdot 2^6 = 2^{12} = 4096$ und damit auch $|D| = 4096$.

War Ihr Ergebnis vielleicht 4095? Dann haben Sie wahrscheinlich einen typischen Denkfehler begangen, um den es bereits in Aufgabe 6.3 ging. Es gilt zwar (siehe Lemma 2.1) $\emptyset \subseteq X$ für jede Menge X , aber das bedeutet *nicht*, dass $\emptyset \in X$ für jede Menge X gilt. Insbesondere gilt in dieser Aufgabe $\emptyset \notin C$ (warum?) und damit $D = C$.

Lösung 15.6. Es gilt $A_4 = \{2, 3\} = A_5$. Damit $|A_q \times A_q| = 100$ gilt, muss A_q nach Lemma 15.1 genau zehn Elemente haben. Und das müssen nach Definition von A_q die ersten zehn Primzahlen sein. Also muss q , weil nach einer Primzahl gefragt wurde, die elfte Primzahl sein: 31.

Lösung 15.7. Überraschenderweise gilt das *nicht*. Die Menge links enthält nur ein Element, nämlich das Paar $(1, (1, 1))$, während die rechte Menge auch nur ein Element enthält, aber ein anderes, nämlich $((1, 1), 1)$. Dieses technische Problem, das nur für die Grundlagenmathematik relevant ist, umgeht man in der Regel, indem man die beiden obigen 2-Tupel mit dem 3-Tupel $(1, 1, 1)$ identifiziert, d. h., man „tut“ einfach so, als wären alle drei Objekte identisch. Wenn man sie gewissermaßen mit der Lupe betrachtet, unterscheiden sie sich, aber rein pragmatisch verhalten sie sich gleich.

Lösung 16.1. Sie können sich das wie ein **Telefonbuch** vorstellen. Für uns Menschen ist es natürlich hilfreich, wenn die Einträge sortiert sind, weil man dadurch wesentlich schneller herausfinden kann, welche Nummer z. B. Frau Müller hat. Aber der Informationsgehalt würde sich nicht ändern, wenn die Reihenfolge durcheinandergewürfelt würde. Irgendwo im Telefonbuch stünde dann immer noch die Nummer von Frau Müller.

Lösung 16.2. Es darf nicht von einem Objekt links mehr als ein Pfeil ausgehen.

Einheitskreis

Lösung 16.3. f_1 ist eine Funktion, denn alle Elemente sind geordnete Paare und f_1 ist rechtseindeutig: Wenn m und n gleich sind, dann sind natürlich auch die Werte $m/3$ und $n/3$ gleich. f_2 beschreibt geometrisch den sogenannten *Einheitskreis*: die Menge aller Punkte (x, y) , die den Abstand 1 vom Punkt $(0, 0)$ haben. Das ist zwar eine Menge von geordneten Paaren, aber keine Funktion: Zu f_2 gehören z. B. die beiden Elemente $(0, 1)$ und $(0, -1)$, also ist f_2 nicht rechtseindeutig. f_3 hat nicht nur geordnete Paare als Elemente, sondern auch die Zahl 3. Daher ist f_3 nicht einmal eine Relation, geschweige denn eine Funktion. f_4 ist eine Funktion. f_5 aber nicht, weil die Rechtseindeutigkeit verletzt ist, denn sowohl $(1, 2)$ als auch $(1, 3)$ gehören zu f_5 . f_6 – die leere Menge – ist eine Funktion. Sie ist erstens eine Menge von Paaren (nämlich von null Paaren) und zweitens ist sie rechtseindeutig, denn eine Relation, die *nicht* rechtseindeutig ist, muss ja mindestens zwei Paare enthalten.

Lösung 16.4. Das sollte hoffentlich kein Problem gewesen sein:

$$\begin{aligned} \text{dom}(f_1) &= \mathbb{N} & \text{rng}(f_1) &= \{n/3 : n \in \mathbb{N}\} \\ \text{dom}(f_2) = \text{rng}(f_2) &= \{x \in \mathbb{R} : |x| \leq 1\} = [-1, 1] \\ \text{dom}(f_4) &= \{1, 3, 4\} & \text{rng}(f_4) &= \{3, 5\} \\ \text{dom}(f_5) &= \{1, 4\} & \text{rng}(f_5) &= \{2, 3, 5\} \\ \text{dom}(f_6) = \text{rng}(f_6) &= \emptyset \end{aligned}$$

Beachten Sie, dass wir Definitions- und Wertebereich nicht nur für Funktionen, sondern allgemeiner für Relationen definiert hatten. Nur für f_3 ergibt die Angabe dieser beiden Mengen keinen Sinn.

Lösung 16.5. Wegen der Rechtseindeutigkeit muss $|f| = |\text{dom}(f)|$ gelten. Außerdem gilt $|f| \geq |\text{rng}(f)|$, aber nicht notwendig $|f| = |\text{rng}(f)|$. Letzteres sieht man bereits am Pfeildiagramm des ersten Beispiels.

Lösung 16.6. Bei komplizierteren Funktionen ist es häufig nicht so einfach, den Wertebereich genau anzugeben. Die Zielmenge sollte man allerdings möglichst „klein“ wählen, um informativ zu sein. Z. B. ist es besser, in (16.2) $\mathbb{Q}$ statt $\mathbb{R}$ zu schreiben, weil man dann sofort sieht, dass auf jeden Fall keine irrationalen Zahlen ausgegeben werden.

Lösung 16.7. So eine Grafik stellt genau dann eine rechtseindeutige Relation dar, wenn es keine vertikale Gerade gibt, die den Graphen mehr als einmal schneidet. Beim Einheitskreis schneidet z. B. die y -Achse den Graphen in den zwei Punkten $(0, -1)$ und $(0, 1)$, was bedeutet, dass diese beiden Paare zur Relation gehören. Daher ist sie nicht rechtseindeutig.

Lösung 16.8. Ich hatte an Funktionen der Form $x \mapsto mx + b$ gedacht, also an die, die man als Geraden visualisieren kann.

Lösung 16.9. Eine Funktion der Form $x \mapsto mx + b$ ist nur dann injektiv, wenn $m \neq 0$ gilt. Ansonsten hat man es nämlich mit einer *konstanten* Funktion der Form $x \mapsto b$ zu tun, die offensichtlich nicht injektiv ist.

Lösung 16.10. Das ist einfach die korrekte Umkehrung der Implikation, um die es im Abschnitt über Aussagenlogik ging.

Lösung 16.11. Eine Funktion ist genau dann injektiv, wenn es keine horizontale Linie gibt, die den Funktionsgraphen mehr als einmal schneidet. Bei der Funktion $x \mapsto x^2$ ist es beispielsweise so, dass eine Gerade parallel zur x -Achse auf der Höhe 2.25 den Funktionsgraphen in den Punkten $(-1.5, 2.25)$ und $(1.5, 2.25)$ schneidet.

Lösung 16.12. f ist nicht injektiv, weil z. B. $f(1) = 1 = f(43)$ gilt. Für g könnte man das Beispiel $g(\{1\}) = 1 = g(\{2\})$ nehmen und $h(2) = 2 = h(11)$ für h . Wahrscheinlich werden Ihre Beispiele anders ausgesehen haben. Die Hauptsache ist, dass Sie das Prinzip verstanden haben.

Lösung 16.13. Das ist natürlich $\text{rng}(f)$, also der Wertebereich von f .

Lösung 16.14. Die Elemente einer Funktion haben allgemein die Form (x, y) und in der Umkehrrelation werden die „umgedrehten“ Paare (y, x) gesammelt. In f gilt $(x, y) = (x, 2x + 8)$ und $(y, x) = (2x + 8, x)$. Wir würden gerne so etwas wie $(y, \dots)$ schreiben, wobei die Punkte für einen Ausdruck stehen, in dem nur noch y vorkommt. Dafür lösen wir $y = 2x + 8$ nach x auf. Das ergibt $y - 8 = 2x$ bzw. $x = y/2 - 4$. Die Umkehrfunktion ist also $f^{-1} = \{(y, y/2 - 4) : y \in \mathbb{R}\}$.

Für g ist die gesuchte Rechenvorschrift $B \mapsto [6] \setminus B = g(B)$, d. h., g^{-1} ist g . Falls Sie das nicht sofort sehen, machen Sie sich vielleicht mit einem Mengendiagramm klar, dass für $A \subseteq [6]$ immer $g(g(A)) = [6] \setminus ([6] \setminus A) = A$ gilt.

Obwohl es sich um dieselbe Rechenvorschrift wie in (16.3) handelt, ist h injektiv, denn der Definitionsbereich enthält keine negativen Zahlen. Die Rechenvorschrift für h^{-1} ist $m \mapsto \sqrt{m}$, wobei zu beachten ist, dass der Definitionsbereich von h^{-1} die Menge $h[\mathbb{N}] = \{n^2 : n \in \mathbb{N}\} = \{0, 1, 4, 9, 16, \dots\}$ ist – siehe Aufgabe 16.13.

Lösung 16.15. Grundsätzlich ist id_A für jede Menge A injektiv und die Umkehrfunktion ist immer id_A . Aber vielleicht hatten Sie auch eine andere Idee.

Lösung 16.16. f ist nicht injektiv, weil u. a. $f(-1) = 1 = f(1)$ gilt.

Lösung 16.17. Nein, denn es gilt beispielsweise $2 \cdot 3 = 3 \cdot 2$, d. h. die beiden unterschiedlichen Tupel $(2, 3)$ und $(3, 2)$ werden auf denselben Funktionswert abgebildet.

Erinnern Sie sich an Aufgabe 15.1. Dort ging es auch um die Multiplikation, die dort A^2 auf B abgebildet hat. Weil die Funktion nicht injektiv ist, hat B weniger Elemente als A^2 .

Lösung 16.18. Nur f_1 ist nicht injektiv: beispielsweise gilt $f_1(0, 0) = 0 = f_1(0, 1)$.

Lösung 16.19. Ein Beispiel wäre die Funktion, die (m, n) auf $2^m 3^n$ abbildet. Dass die injektiv ist, folgt aus dem [Fundamentalsatz der Arithmetik](#). So eine „Codierung“ von Tupeln nennt man auch [Gödelisierung](#). Sie ist ein wichtiges Werkzeug in der [theoretischen Informatik](#).

Lösung 16.20. Das würde dazu führen, dass der Fall $x = 0$ doppelt auftaucht. Es würde in diesem Fall nichts ändern, weil in beiden Fällen der Funktionswert derselbe wäre, aber man sollte so etwas trotzdem möglichst vermeiden, weil man damit seine Leser verwirren könnte.

Lösung 16.21. Sowohl bei f_1 als auch bei f_2 ist die Definition unglücklich, weil die natürlichen Zahlen von 4 bis 9 jeweils beide Bedingungen erfüllen. Bei f_1 ist das aber nicht so schlimm, weil sich in beiden Fällen für alle sechs Zahlen der Funktionswert null ergibt. f_2 ist jedoch keine Funktion, denn es ist beispielsweise nicht klar, ob dem Argument 7 der Funktionswert 0 oder der Funktionswert 1 zugeordnet werden soll.

Lösung 16.22. $f = \{(1, 5), (2, 2), (3, 3), (4, 4), (5, 1), (6, 6)\}$. $\text{rng}(f)$ ist $[6]$.

Lösung 16.23. Es kommt $\{(1, 5), (2, 2), (3, 3), (5, 5)\}$ heraus und das ist nicht dasselbe wie $f \circ g$. Die Komposition von Funktionen ist also im Allgemeinen *nicht* kommutativ.

Lösung 16.24. Wir haben $(f \circ g)(x) = (2^x)^2 = 2^{2x}$ und $(g \circ f) = 2^{x^2}$. (Siehe Lemma 5.8.) Dass die beiden Rechenvorschriften, die wir erhalten haben, unterschiedlich *aussehen*, bedeutet noch nicht zwangsläufig, dass $f \circ g$ und $g \circ f$ nicht gleich sind. (Vielleicht kennt ja jemand eine geschickte Umformung, die aus der einen Rechenvorschrift die andere macht ...) Daher brauchen wir, siehe Aufgabenstellung, ein konkretes Gegenbeispiel; so etwas wie $(f \circ g)(1) = 4$ und $(g \circ f)(1) = 2$.

Übrigens hätten Sie in diesem Fall auch Pech haben und zufällig den „falschen“ Wert einsetzen können: $(f \circ g)(2) = 16 = (g \circ f)(2)$. Ihnen ist natürlich klar, dass das *nicht* bedeutet, dass die beiden Funktionen identisch sind, oder?

Lösung 16.25. Wir bekommen

$$f \circ g: \begin{cases} \mathbb{N} \rightarrow \mathbb{N} \\ x \mapsto x \end{cases} \quad \text{und} \quad g \circ f: \begin{cases} \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N} \times \mathbb{N} \\ (x, y) \mapsto (x, x) \end{cases}$$

heraus. $f \circ g$ ist also einfach $\text{id}_{\mathbb{N}}$.

Lösung 16.26. Ein Beispiel wären die durch $f(x) = x$ und $g(x) = x + 1$ auf $\mathbb{R}$ definierten Funktionen.

Lösung 16.27. Definieren wir $f(n) = \lfloor n/2 \rfloor$ und $g(n) = 2n$, dann ergibt sich $\text{id}_{\mathbb{N}}$ als Komposition $f \circ g$. Man erhält also eine injektive Funktion, obwohl f ersichtlich nicht injektiv ist. Das liegt anschaulich formuliert daran, dass g nicht alle natürlichen Zahlen „trifft“: f hat für jede gerade Zahl und die direkt auf sie folgenden denselben Funktionswert. Aber durch den „Umweg“ über g werden in f nur gerade Zahlen eingesetzt.

Dass man für (ii) kein Beispiel finden kann, kann man sich folgendermaßen klar machen: Ist g nicht injektiv, so gibt es natürliche Zahlen m und n mit $m \neq n$ und $g(m) = g(n)$. Setzt man nun in f ein, dann hat man $f(g(m)) = f(g(n))$.

Lösung 16.28. Ein sehr simples Beispiel wäre $f = \{(1, 0), (2, 0)\}$ und $g = \{(0, 0), (1, 0)\}$. Dann ist $f \circ g$ die leere Menge, also injektiv. Falls Ihnen das zu „banal“ ist, dann fügen Sie sowohl zu f als auch zu g noch das Paar $(3, 3)$ hinzu. $f \circ g$ wäre dann $\{(3, 3)\}$ und nach wie vor injektiv.

Warum widerspricht das nicht Aufgabe 16.27? Dort wurde in der Aufgabenstellung festgelegt, dass der Wertebereich von g eine Teilmenge des Definitionsbereichs von f sein musste. Hier ist es aber so, dass g die Ausgabe 0 liefert, die man gar nicht als Eingabe in f einsetzen kann.

Lösung 16.29. Mit $f(0, 0) = 0 = f(0, 1)$ sieht man, dass f nicht injektiv ist. Und mit $(g \circ f)(0, 0) = (0, 0) = (g \circ f)(0, 1)$ sieht man, dass $g \circ f$ ebenfalls nicht injektiv ist.²⁶ Die anderen beiden Funktionen sind offenbar injektiv.

Lösung 16.30. Mit $f(0, 0) = 1 = f(1, 0)$ sieht man, dass f nicht injektiv ist. Dieses Gegenbeispiel muss dann auch für $g \circ f$ funktionieren:

$$(g \circ f)(0, 0) = (0, 2) = (g \circ f)(1, 0)$$

g ist offensichtlich injektiv. Die Rechenvorschrift für $f \circ g$ ist

$$(f \circ g)(x) = f(g(x)) = f(x - 1, x + 1) = x + 2$$

und das beschreibt sicherlich auch eine injektive Funktion.

²⁶Wie in Aufgabe 16.27 folgt die Nicht-Injektivität von $g \circ f$ aus der von f und man kann dieselben Werte einsetzen, um ein Gegenbeispiel zu bekommen.

Lösung 16.31. Alle vier sind injektiv. Man kommt evtl. bereits durch das Einsetzen von ein paar Werten und „scharfes Hinsehen“ darauf, dass f und g wohl injektiv sind. Man kann das aber auch ausführlich begründen, wenn man sich daran erinnert (siehe Seite 129), dass Injektivität bedeutet, dass man ein Verfahren angeben kann, das zu jeder Ausgabe eindeutig die Eingabe rekonstruieren kann. Für g erhält man die Eingabe zur Ausgabe (u, v) durch $(v - 1, u + 1)$. Mit anderen Worten: g ist seine eigene Umkehrfunktion. Für f erhält man die Eingabe zur Ausgabe (u, v) durch $((u + v)/2, (u - v)/2)$. (Das läuft auf das aus der Schule bekannte Lösen eines ganz einfachen **linearen Gleichungssystems** hinaus: $u = x + y$ und $v = x - y$ werden nach x und y aufgelöst.)

Die Funktionen $f \circ g$ und $g \circ g$ müssen dann als Kompositionen injektiver Funktionen zwangsläufig auch injektiv sein. Der Vollständigkeit halber geben wir noch die Rechenvorschrift von $f \circ g$ an:

$$(f \circ g)(x, y) = f(g(x, y)) = f(y - 1, x + 1) = (x + y, y - x - 2).$$

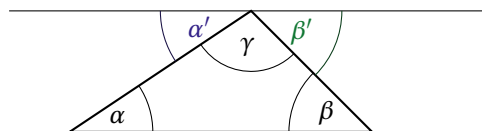
Dass $g \circ g = \text{id}_{\mathbb{Z} \times \mathbb{Z}}$ gilt, wurde im letzten Absatz bereits erklärt.

Lösung 17.1. $\alpha = 20^\circ$, $\beta = 340^\circ$, $\gamma = 320^\circ$, $\delta = 40.3^\circ$.

Ihnen ist sicher aufgefallen, dass das eigentlich genau wie **modulare Arithmetik** modulo 360 funktioniert. Allerdings mit dem kleinen Unterschied (siehe Beispiel δ), dass man es auf alle reellen Zahlen und nicht nur auf ganze anwendet.

Lösung 17.2. Die drei anderen spitzen Winkel haben ebenfalls 22 Grad, während die vier stumpfen Winkel 158 Grad haben.

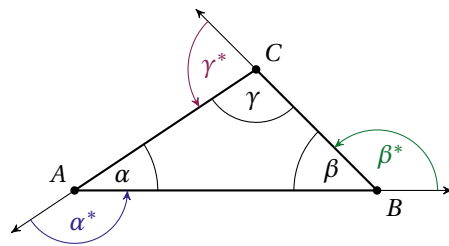
Lösung 17.3. Wir zeichnen zwei weitere Winkel in die Skizze ein:



Offensichtlich sind α und α' sowie β und β' Wechselwinkel und daher jeweils gleich groß. Zudem sieht man sofort, dass α' , β' und γ sich zu 180° ergänzen. Somit gilt dann natürlich auch

$$\alpha + \beta + \gamma = \alpha' + \beta' + \gamma = 180^\circ.$$

Es gibt natürlich noch andere Möglichkeiten, sich diese Tatsache klarzumachen. Hier eine Alternative: Stellen Sie sich eine Ameise vor, die einmal um das Dreieck herum läuft. Am Anfang befindet sie sich wie auf dem Bild unten auf der Strecke von A nach B und schaut in Richtung der Ecke B .

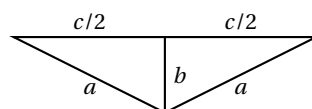


Am Punkt B muss sich die Ameise um den Winkel β^* drehen, dann geht sie weiter in Richtung C . Bei C dreht sie sich um γ^* und so weiter. Am Ende ihrer Reise ist sie wieder am Ausgangspunkt und schaut in dieselbe Richtung wie am Anfang. Daher muss sie sich insgesamt um 360° gedreht haben. Andererseits sind β und β^* Ergänzungswinkel, d. h. $\beta + \beta^* = 180^\circ$, und das gilt ebenso für die anderen beiden Winkelpaare. Die Gesamtdrehung setzt sich also so zusammen:

$$\begin{aligned} 360^\circ &= \alpha^* + \beta^* + \gamma^* = (180^\circ - \alpha) + (180^\circ - \beta) + (180^\circ - \gamma) \\ &= 540^\circ - (\alpha + \beta + \gamma) \end{aligned}$$

Aus dieser Beziehung folgt offenbar sofort $\alpha + \beta + \gamma = 180^\circ$.

Lösung 17.5. Wir verbinden den Mittelpunkt der Seite c mit dem gegenüberliegenden Eckpunkt des Dreiecks.

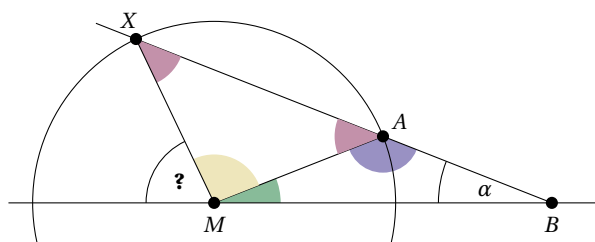


Dadurch haben wir das Dreieck in zwei Dreiecke zerlegt, die drei paarweise gleich lange Seiten – a , $c/2$ und b – haben. Nach dem Kongruenzsatz SSS sind die beiden Dreiecke kongruent und daher müssen auch die entsprechenden Winkel paarweise gleich sein, insbesondere die beiden, um die es uns geht (und die nun jeweils gegenüber der Seite b liegen).

Nebenbei sieht man aufgrund dieser Winkelgleichheit auch, dass die Seite b den Winkel gegenüber von c halbiert und dass die beiden Dreiecke rechtwinklig sind.

Lösung 17.6. Weil ein gleichseitiges Dreieck insbesondere auch ein gleichschenkeliges ist, kann man Aufgabe 17.5 anwenden und folgern, dass zwei der drei Winkel gleich sein müssen. Das kann man aber für jedes Paar von Seiten bzw. Winkeln machen und daher müssen die Winkel alle gleich groß sein. Weil ihre Summe nach Aufgabe 17.3 180° beträgt, muss jeder von ihnen das Maß 60° haben.

Lösung 17.7. Die folgende Skizze illustriert die Lösung.



Weil MBA nach Aufgabenstellung ein gleichschenkliges Dreieck ist, sind der Winkel α und der grüne Winkel bei M gleich groß. (Siehe Aufgabe 17.5.) Der blaue Winkel bei A beträgt daher $180^\circ - 2\alpha$. Der rote Winkel bei A ist der Nebenwinkel dazu und hat somit das Maß 2α . Da ferner sowohl $[MA]$ als auch $[MX]$ Radien des Kreises sind, ist das Dreieck MAX ebenfalls gleichschenkl. Deshalb ist der rote Winkel bei X so groß wie der rote bei A . Für den gelben Winkel bei M ergibt sich somit $180^\circ - 4\alpha$. Der gesuchte Winkel bildet zusammen mit den anderen beiden Winkeln bei M 180 Grad. Damit ist sein Maß $180^\circ - (180^\circ - 4\alpha) - \alpha = 3\alpha$.

Lösung 17.8. Das Verhältnis von x_3 zu x_1 ist $x_3/x_1 = 1.5$ und damit dasselbe wie das Verhältnis $x_4/x_2 = 1.5$ von x_4 zu x_2 . Man hätte aber auch das Verhältnis von x_1 zu x_3 (also den Kehrwert von 1.5) mit dem von x_2 zu x_4 vergleichen können. Ebenso entspricht das Verhältnis von x_2 zu x_1 dem von x_4 zu x_3 . Das ergibt sich aus dem bereits berechneten Verhältnissen durch

$$\frac{x_4}{x_3} = \frac{1.5x_2}{1.5x_1} = \frac{x_2}{x_1}.$$

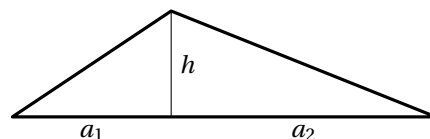
Das heißt aber nicht, dass Sie irgendwelche Paare nehmen können. Das Verhältnis von x_3 zu x_2 ist beispielsweise *nicht* dasselbe wie das von x_4 zu x_1 !

Lösung 17.9. Setzen wir $x_1 = |ZA|$, $x_2 = |AA'|$, $y_1 = |ZB|$ und $y_2 = |BB'|$, so sind offenbar alle vier Zahlen positiv und es gilt $|ZA'| = x_1 + x_2$ sowie $|ZB'| = y_1 + y_2$. Die erste Aussage besagt dann $x_1/(x_1 + x_2) = y_1/(y_1 + y_2)$. Bildet man die Kehrwerte, so erhält man $1 + x_2/x_1 = (x_1 + x_2)/x_1 = (y_1 + y_2)/y_1 = 1 + y_2/y_1$. Subtrahiert man nun auf beiden Seiten 1 und bildet erneut die Kehrwerte, hat man die zweite Aussage. Durch die umgekehrten Umformungen zeigt man, dass aus der zweiten Gleichung die erste folgt. Die Äquivalenz zur dritten Aussage erhält man mit einer analogen Rechnung.

Lösung 17.10. Wenn in zwei Dreiecken alle Winkel gleich sind, sind sie – wie wir gerade gelernt haben – ähnlich. Aber sie sind nicht notwendig kongruent. Die Skizzen vor dieser Aufgabe zeigen Beispiele für Dreiecke unterschiedlicher Größe, die lauter identische Winkel haben.

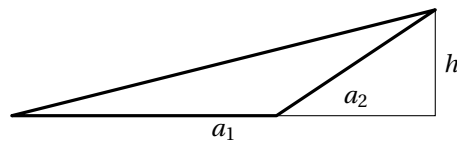
Lösung 17.11. Nach Pythagoras gilt $10^2 = 8^2 + b^2$. Es folgt $b^2 = 10^2 - 8^2 = 100 - 64 = 36$, also $b = 6$.

Lösung 17.12. Wir nennen die waagerechte Kathete des linken rechtwinkligen Dreiecks a_1 und die des rechten a_2 . Die senkrechte Kathete hat bei beiden Dreiecken dieselbe Länge h .



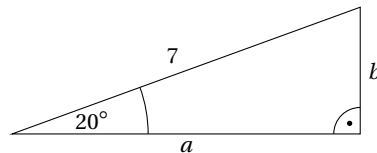
Somit hat das linke Dreieck die Fläche $a_1 h/2$ und für das rechte erhalten wir $a_2 h/2$. Als Fläche des großen Dreiecks bekommen wir daher $(a_1 + a_2)h/2$. Dabei ist $a_1 + a_2$

die Länge der sogenannten *Grundseite* und h die *Höhe* auf der Grundseite. Und das funktioniert auch dann noch, wenn die Höhe außerhalb des Dreiecks liegt:



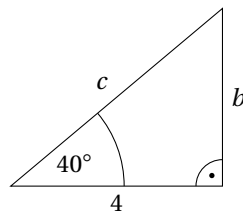
Hier ist a_1 die Kathete eines rechtwinkligen Dreiecks, das größer als das Dreieck ist, dessen Fläche wir wissen wollen. Um die gesuchte Fläche zu erhalten, müssen wir von der Fläche des großen rechtwinkligen Dreiecks die des kleinen (mit der Kathete a_2) subtrahieren. Das ergibt $(a_1 - a_2)h/2$ und das ist korrekt, weil $a_1 - a_2$ offenbar die Länge der Grundseite ist.

Lösung 17.13. Das Dreieck sollte ungefähr so aussehen:



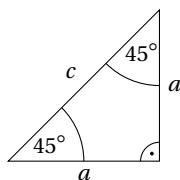
Mit $b/7 = \sin 20^\circ \approx 0.342$ ergibt sich $b \approx 7 \cdot 0.342 = 2.394$. Ebenso erhält man für die andere Kathete $a = 7 \cos 20^\circ \approx 7 \cdot 0.940 = 6.580$.

Lösung 17.14. Das Dreieck sieht in etwa folgendermaßen aus:



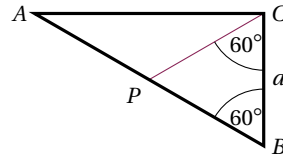
Wegen $b/4 = \tan 40^\circ \approx 0.839$ erhalten wir $b \approx 4 \cdot 0.839 = 3.356$. Mit Pythagoras folgt nun $c = \sqrt{4^2 + b^2} \approx 5.221$.

Lösung 17.15. Da die Winkelsumme im Dreieck 180° beträgt, müssen *beide* nicht-rechten Winkel das Maß 45° haben, wenn einer von beiden dieses Maß hat. Daher müssen natürlich auch die beiden Katheten dieselbe Länge haben. Als Tangens ergibt sich damit sofort $a/a = 1$.



Mit Pythagoras erhält man für die Länge der Hypotenuse c nun $\sqrt{a^2 + a^2} = a\sqrt{2}$. Daher sind die Werte für den Sinus und den Kosinus von 45° beide $a/c = 1/\sqrt{2}$.

Lösung 17.16. Damit sich insgesamt 180 Grad ergeben, muss der dritte Winkel im Dreieck BCP ebenfalls das Maß 60° haben. Damit ist dieses Dreieck gleichseitig. Nennen wir die waagerechte Kathete a , so haben also sowohl die rote Strecke $[PC]$ als auch $[BP]$ die Länge a .

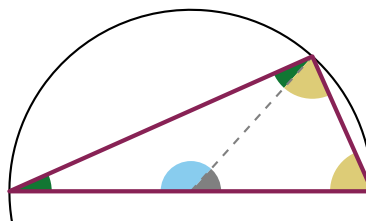


Der Winkel des Dreiecks CAP an der Ecke C beträgt 30° , weil bei C ja der rechte Winkel des Dreiecks ABC liegt. Außerdem ergibt sich (wieder wegen der Winkelsumme im Dreieck), dass der Winkel bei A ebenfalls das Maß 30° hat. Damit ist APC ein gleichschenkliges Dreieck und wir können folgern, dass $[PA]$ dieselbe Länge wie $[PC]$, also a , hat. Die Hypotenuse des rechtwinkligen Dreiecks hat somit die Länge $2a$ und als Kosinus von 60° ergibt sich $a/(2a) = 1/2$. Jetzt können wir mit Pythagoras die Länge der zweiten Kathete ausrechnen und erhalten $\sqrt{(2a)^2 - a^2} = a\sqrt{3}$. Daher hat der Sinus von 60° den Wert $a\sqrt{3}/(2a) = \sqrt{3}/2$ und für den Tangens erhalten wir $a\sqrt{3}/a = \sqrt{3}$.

Lösung 17.17. Wir können das Dreieck von oben benutzen und müssen nur die Rollen von An- und Gegenkathete vertauschen. Damit erhalten wir:

$$\sin 30^\circ = 1/2, \quad \cos 30^\circ = \sqrt{3}/2 \quad \text{und} \quad \tan 30^\circ = 1/\sqrt{3}.$$

Lösung 17.18. Der gestrichelte Radius teilt das rote Dreieck in zwei gleichschenklige Dreiecke auf. (Die beiden gleich langen Schenkel sind jeweils Radien des Kreises.)



Nennen wir den grauen Winkel im rechten Dreieck α , so haben die beiden gelben Winkel jeweils das Maß $(180^\circ - \alpha)/2 = 90^\circ - \alpha/2$. Der blaue Winkel im linken Dreieck beträgt $180^\circ - \alpha$ und deshalb müssen die beiden grünen Winkel jeweils das Maß $(180^\circ - (180^\circ - \alpha))/2 = \alpha/2$ haben. Addiert man einen grünen und einen gelben Winkel, so ergeben sich genau 90 Grad.

Lösung 17.19. Ein Fünfeck kann man in drei Dreiecke zerlegen. Da die Summe der Innenwinkel des Dreiecks 180° beträgt, ergibt sich für die entsprechende Summe im Fünfeck $3 \cdot 180^\circ = 540^\circ$.

Lösung 17.20. Mit jeder weiteren Ecke kann man ein weiteres Dreieck abschneiden. Bei fünf Ecken ergibt sich nach der letzten Aufgabe als Summe $(5 - 2) \cdot 180^\circ$,

beim Sechseck dann $(6 - 2) \cdot 180^\circ$ und so weiter. Allgemein ist die Summe der Innenwinkel in einem Polygon mit n Ecken also $(n - 2) \cdot 180^\circ$. (Natürlich ergibt diese Ausgabe nur dann Sinn, wenn die Seiten des Polygons sich nicht überschneiden.)

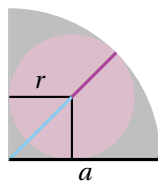
Lösung 17.21. Wir betrachten ein Viertel des Quadrats:



Wenn wir die Seitenlänge des ursprünglichen Quadrats a nennen, dann hat dieses kleine Quadrat die Seitenlänge $a/2$. Die orangefarbene Diagonale hat nach Pythagoras dann die Länge $a/\sqrt{2}$. (Rechnen Sie nach!) Da der Durchmesser des grauen Kreises sicher ebenfalls $a/2$ ist und der Radius des roten Kreises offenbar die Hälfte dessen ausmacht, was übrigbleibt, wenn man von der Diagonale den Durchmesser entfernt, erhält man diese Lösung:

$$\frac{1}{2} \cdot \left(\frac{a}{\sqrt{2}} - \frac{a}{2} \right) = \frac{\sqrt{2} - 1}{4} \cdot a$$

Lösung 17.22. Für die erste Teilaufgabe schauen wir uns das obere rechte Viertel an und nennen den Radius des grauen Kreises a .



Nennen wir den gesuchten Radius des roten Kreises r , so ist die Länge des blauen Streckenstücks in der obigen Skizze nach Pythagoras $\sqrt{2}r$ und die Länge des violetten Stücks r . Die beiden Strecken zusammen ergeben offenbar a :

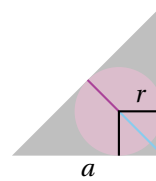
$$a = r + \sqrt{2}r = (1 + \sqrt{2})r$$

Somit ergibt sich:²⁷

$$r = \frac{a}{1 + \sqrt{2}} = \frac{a}{1 + \sqrt{2}} \cdot \frac{1 - \sqrt{2}}{1 - \sqrt{2}} = a(\sqrt{2} - 1)$$

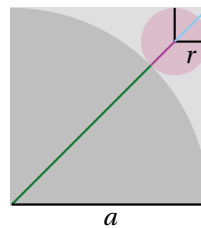
Wenn man das Prinzip der obigen Lösung verstanden hat, dann stellt man fest, dass sich die anderen beiden Teilaufgaben sehr ähnlich lösen lassen. Für den zweiten Teil nennen wir die Seitenlänge des Quadrats a und den gesuchten Radius wieder r .

²⁷Man entfernt die Wurzel aus dem Nenner (und vereinfacht so den Bruch), indem man so erweitert, dass man die dritte binomische Formel anwenden kann. Diesem Trick kennen Sie schon von den komplexen Zahlen.



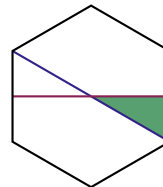
Hier setzen wir die Hälfte der Diagonalen des Quadrats aus zwei Teilen zusammen:
 $a/\sqrt{2} = r + \sqrt{2}r = (1 + \sqrt{2})r$. Auflösen nach r liefert $r = a/(2 + \sqrt{2})$.

Im dritten Fall benennen wir wie folgt:



Wir erhalten $\sqrt{2}a = \sqrt{2}r + r + a$ bzw. $a(\sqrt{2} - 1) = r(\sqrt{2} + 1)$. Löst man nach r auf und vereinfacht, so ergibt sich $r = a(3 - 2\sqrt{2})$.

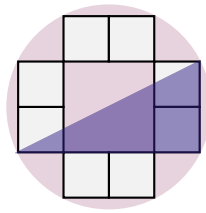
Lösung 17.23. Wie vor der Aufgabe erwähnt, empfiehlt es sich, ein rechtwinkliges Dreieck zu finden, von dem man schon ein paar Eigenschaften kennt. Ein solches wurde hier grün eingezeichnet:



Wir wissen nach Aufgabenstellung, dass die (blaue) Hypotenuse die Länge 2 hat. Außerdem muss der Winkel links das Maß $360^\circ/12 = 30^\circ$ haben, da man offenbar zwölf solche Dreiecke im Sechseck unterbringen könnte und da dieses regelmäßig ist. Die rote Kathete hat daher die Länge $2 \cdot \cos 30^\circ = 2 \cdot \sqrt{3}/2 = \sqrt{3}$. Die Antwort ist somit $2\sqrt{3}$.

Alternativ kann man argumentieren, dass man sechs gleichseitige Dreiecke erhält, wenn man den Mittelpunkt des Sechsecks mit den sechs Ecken verbindet. Daher muss die Länge der Seiten des Sechsecks die Hälfte von 4 (Länge der blauen Linie) sein. Man kennt dann zwei Seiten des grünen Dreiecks (außer der Hypotenuse noch die schwarze Kathete mit der Länge 1) und kann die Dritte mit Pythagoras und ohne Kenntnis der Winkel berechnen.

Lösung 17.24. Das rechtwinklige Dreieck, das uns in diesem Fall hilft, wurde hier blau eingezeichnet:



Da der Kreis den Durchmesser 2 hat, ist das die Länge der Hypotenuse des blauen Dreiecks. Nennt man die Seitenlänge eines der Quadrate a , so ergibt sich mit Pythagoras $(2a)^2 + (4a)^2 = 2^2$ und damit $a = 1/\sqrt{5}$. Die Fläche eines Quadrats ist also $1/5$ und die aller acht Quadrate daher $8/5$.

Lösung 17.25. Der Winkel zwischen Diagonale und Seitenlänge des Quadrats beträgt natürlich 45° . Damit folgt $\alpha = 15^\circ$. Nach Pythagoras ist die Seitenlänge des Quadrats $35/\sqrt{2}$. Das ist gleichzeitig die Länge der Ankathete an den Winkel 2α im unteren rechtwinkligen Dreieck. Dessen Hypotenuse hat daher die Länge

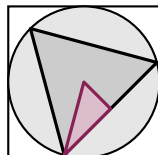
$$c = \frac{35/\sqrt{2}}{\cos 2\alpha} = \frac{35}{\sqrt{2} \cdot \cos 30^\circ}.$$

In dem oberen rechtwinkligen Dreieck ist die rote Seite die Gegenkathete zu α . Die gesuchte Länge ist demnach

$$c \cdot \tan \alpha = \frac{35 \cdot \tan 15^\circ}{\sqrt{2} \cdot \cos 30^\circ}. \quad (\text{L-2})$$

Man könnte nun noch den exakten Wert für $\cos 30^\circ$ aus Aufgabe 17.17 einsetzen und etwas vereinfachen, aber für diese Aufgabe reicht der Ausdruck (L-2).

Lösung 17.26. Die Ecken des (offensichtlich rechtwinkligen) roten Dreiecks in der folgenden Skizze sind der Mittelpunkt des Kreises, der Mittelpunkt einer Seite des grauen Dreiecks und eine Ecke des grauen Dreiecks.



Da das Quadrat die Fläche 75 hat, ist seine Seitenlänge $\sqrt{75}$ und das ist dann auch der Durchmesser des Kreises. Die Hypotenuse des roten Dreiecks hat daher die Länge $\sqrt{75}/2$. Weil das graue Dreieck gleichseitig ist, haben seine Innenwinkel alle das Maß 60° . Der untere Winkel des roten Dreiecks beträgt daher 30° und damit ergibt sich für die Länge der Ankathete zu diesem Winkel

$$\frac{\sqrt{75}}{2} \cdot \cos 30^\circ = \frac{\sqrt{75}}{2} \cdot \frac{\sqrt{3}}{2} = \frac{15}{4}.$$

Die gesuchte Seitenlänge ist somit $15/2$.

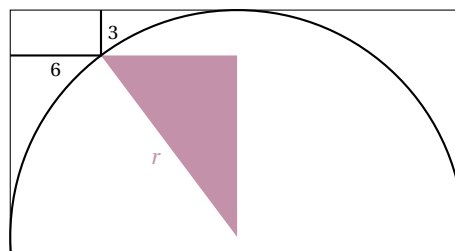
Es ist übrigens immer gut, zu überprüfen, ob das vorgebliche Ergebnis nicht eventuell eine völlig falsche Größenordnung hat. Hier sieht man z. B. an der Skizze,

dass die gesuchte Seitenlänge ein bisschen kürzer als die des Quadrats ist. $\sqrt{75}$ liegt zwischen $8 = \sqrt{64}$ und $9 = \sqrt{81}$, daher scheint das Resultat $15/2 = 7.5$ durchaus realistisch zu sein. Hätten Sie etwa 1.3 (viel zu klein) oder 512 (viel zu groß) herausbekommen, so wäre hingegen klar gewesen, dass das nicht richtig sein kann.

Lösung 17.27. Wenn man das Dreieck halbiert, erhält man zwei identische (gespiegelte) rechtwinklige Dreiecke mit einer Hypotenuse der Länge 13 und einer Kathete der Länge 5. (Wenn Ihnen das nicht klar ist, dann machen Sie sich eine Skizze.) Nach dem [Satz des Pythagoras](#) hat die andere Kathete die Länge $\sqrt{13^2 - 5^2} = 12$. Die gesuchte Fläche ist daher $5 \cdot 12 = 60$.

Lösung 17.28. Nach dem [Satz des Thales](#) ist der Winkel bei C ein rechter. Daher müssen die anderen beiden Winkel zusammen das Maß $180^\circ - 90^\circ$ haben, woraus $\alpha = 18^\circ$ folgt. Die Hypotenuse des rechtwinkligen Dreiecks ABC hat nach Aufgabenstellung die Länge 100. Die Länge von $[BC]$ ist daher $100 \cdot \sin 2\alpha = 100 \cdot \sin 36^\circ$ oder alternativ $100 \cdot \cos 54^\circ$.

Lösung 17.29. Man betrachte das rote rechtwinklige Dreieck, dessen untere rechte Ecke der Mittelpunkt des Kreises ist.



Nennt man den Radius r , so liefert Pythagoras für dieses Dreieck den Zusammenhang $r^2 = (r - 6)^2 + (r - 3)^2$, der sich zu $r^2 - 18r + 45 = 0$ umformen lässt. Diese [quadratische Gleichung](#) hat die beiden Lösungen 3 und 15, von denen offenbar nur die zweite für die Aufgabe infrage kommt.

Lösung 18.1. Nach [Satz 18.1](#) hat der Kreisumfang (der einer vollen Umdrehung entspricht) die Länge 2π . Die Hälfte davon ist π .

Lösung 18.2. Der Kosinus verschwindet, wenn der „Zeiger“ im Kreis senkrecht nach oben oder unten zeigt, also bei $\pi/2$ und $3\pi/2$ bzw. ganz allgemein für alle Werte der Form $k\pi + \pi/2$, wobei k irgendeine ganze Zahl ist. Analog ergibt sich, dass der Sinus für Werte der Form $k\pi$ verschwindet.

Lösung 18.3. Nach dem [Satz des Pythagoras](#) ist die Länge der Hypotenuse des rechtwinkligen Dreiecks 17. Der Ergänzungswinkel α' zu α ist ein Innenwinkel dieses Dreiecks und nach Definition gilt $\sin \alpha' = 8/17$. Wegen $\alpha = \pi - \alpha'$ folgt mit [Lemma 18.2](#)

$$\sin \alpha = \sin(\pi - \alpha') = -\sin(\alpha' - \pi) = \sin \alpha' = 8/17.$$

Lösung 18.4. Nach dem **Satz des Pythagoras** ist die Länge der Hypotenuse $c = \sqrt{k_1^2 + k_2^2}$ und damit folgt $\sin \alpha = k_1 / c$. Quadriert man das, so erhält man $\sin^2 \alpha = k_1^2 / (k_1^2 + k_2^2)$.

Lösung 18.5. So zum Beispiel:

$$\cos(2\pi/3) = \cos(\pi/2 + \pi/6) = -\cos(\pi/2 - \pi/6) = -\cos(\pi/3) = -1/2$$

$$\cos(5\pi/4) = \cos(\pi + \pi/4) = -\cos(\pi/4) = -1/\sqrt{2}$$

$$\cos(-\pi/3) = \cos(\pi/3) = 1/2$$

$$\cos(8\pi + 1) = \cos 1 \approx 0.54$$

$$\sin(2\pi/5) = \sin(-\pi/10 + \pi/2) = \cos(-\pi/10) = \cos(\pi/10) \approx 0.95$$

$$\sin(7\pi/4) = \sin(5\pi/4 + \pi/2) = \cos(5\pi/4) = -\cos(\pi/4) = -1/\sqrt{2}$$

$$\cos(20^\circ) = \cos(20\pi/180) = \cos(\pi/9) \approx 0.94$$

$$\begin{aligned} \sin(250^\circ) &= \sin(250\pi/180) = \sin(25\pi/18) = \sin(7\pi/18 + \pi) \\ &= -\sin(7\pi/18) = -\sin(-2\pi/18 + \pi/2) = -\cos(2\pi/18) \\ &= -\cos(\pi/9) \approx -0.94 \end{aligned}$$

Die Zahlenwerte, die Ihr Taschenrechner ausgeben würde, stehen dort nur zur Information. Nach denen war ja nicht gefragt. (Aber man kann mit einem funktionierenden Taschenrechner seine Ergebnisse selbst überprüfen.)

Die Geschichte mit dem Taschenrechner ist zwar Unsinn, aber die Aufgabe hat durchaus praktische Relevanz. Computer beschränken sich in der Regel wirklich darauf, Kosinus oder Sinus nur für ein kleines Intervall wirklich auszurechnen, und berechnen die Werte für andere Winkel so ähnlich wie hier.

Lösung 18.6. Das erste Additionstheorem liefert

$$\cos 2x = \cos(x+x) = \cos^2 x - \sin^2 x.$$

Die aus dem Satz des Pythagoras hergeleitete Formel $\sin^2 x + \cos^2 x = 1$ ist äquivalent zu $\sin^2 x = 1 - \cos^2 x$. Setzt man das ein, so hat man $\cos 2x = 2\cos^2 x - 1$. Nach $\cos x$ aufgelöst erhält man $\cos x = \pm\sqrt{(1+\cos 2x)/2}$. Man beachte, dass das **Vorzeichen nicht eindeutig bestimmt** ist. Ob der Kosinus positiv oder negativ ist, hängt davon ab, in welchem Bereich sich x befindet – siehe Lemma 18.2.

Lösung 18.7. Nach der besagten Lösung und mit der Tabelle vor der Aufgabe ergibt sich

$$\begin{aligned} \cos(\pi/8) &= \sqrt{(1+\cos(\pi/4))/2} = \sqrt{(1+1/\sqrt{2})/2} \\ &= \sqrt{(2+\sqrt{2})/4} = \sqrt{2+\sqrt{2}}/2 \approx 0.924 \end{aligned}$$

(Den Näherungswert 0.924 sollten Sie natürlich nicht ohne Taschenrechner ermitteln.) Der Wert ist positiv, weil $\pi/8$ in einem Bereich liegt, in dem der Kosinus positive Funktionswerte hat.

Diese Aufgabe war nur als Fingerübung im Rechnen gedacht. Aber ungefähr so musste man in den vergangenen Jahrhunderten tatsächlich vorgehen. Im Video *Wie hat man früher Sinus und Kosinus berechnet?* erfahren Sie mehr darüber.

Lösung 18.8. Die Antwort auf die erste Frage ist *ja*, man kann z. B. $a = b = 1$ und $c = 3$ nehmen oder ganz allgemein Zahlen, so dass die Summe zweier dieser Zahlen kleiner als die dritte ist. Daraus wird kein Dreieck.

Die Antwort auf die zweite Frage ist jedoch *nein*. Gilt $a^2 + b^2 = c^2$, so ist c offenbar die größte der drei Zahlen. Es kann dann nicht $a + b < c$ gelten, denn daraus würde $a^2 + b^2 < a^2 + 2ab + b^2 = (a + b)^2 < c^2$ folgen.

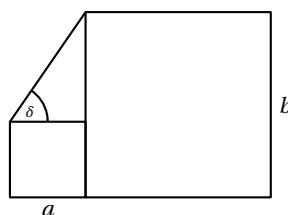
Lösung 18.9. Nennt man den Abstand c , so gilt $c^2 = 4^2 + 6^2 - 2 \cdot 4 \cdot 6 \cos(23^\circ)$ nach dem Kosinussatz. Das ergibt als Näherungswert $c \approx 2.80$.

Lösung 19.1. Das ist natürlich $\sqrt{x^2 + y^2}$ nach dem **Satz des Pythagoras**.

Lösung 19.3. Der Arkussinus liefert nur Werte zwischen $-\pi/2$ und $\pi/2$, so dass die „naive“ Formel $\alpha = \arcsin(y/r)$ nur für Punkte rechts von der y -Achse richtig wäre. Man kann es z. B. so machen:

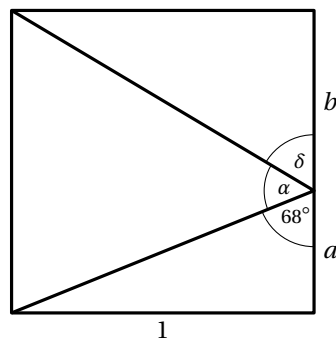
$$\alpha = \begin{cases} \arcsin(y/r) & x \geq 0 \\ \pi - \arcsin(y/r) & x < 0 \end{cases}$$

Lösung 19.4. Bezeichnen wir die Seiten der Quadrate mit a und b und den gesuchten Winkel mit δ , so muss wegen des Flächenverhältnisses $b = a\sqrt{6}$ gelten.



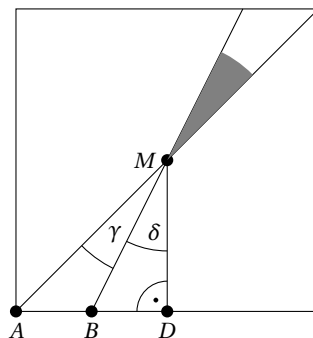
Ferner haben Gegen- und Ankathete von δ in dem rechtwinkligen Dreieck links oben die Längen $b - a$ und a . Ihr Quotient (der Tangens von δ) ist also $(b - a)/a = \sqrt{6} - 1$. Berechnet man den Arkustangens dieser Zahl, so erhält man den Winkel. (Es sind etwas mehr als 55 Grad.)

Lösung 19.5. Die Lösung hängt offenbar nicht vom Maßstab ab. O. B. d. A. können wir also als Seitenlänge des Quadrats 1 annehmen. (Das macht das Rechnen leichter.)



Mit den Bezeichnungen aus der Skizze gilt $\tan(68^\circ) = 1/a$ im unteren rechtwinkligen Dreieck. Damit berechnet man a und kennt dann auch $b = 1 - a$. Im oberen rechtwinkligen Dreieck gilt nun $\tan \delta = 1/b$. Mit dem Arkustangens erhält man so δ . α bekommt man schließlich, weil die drei Winkel zusammen 180 Grad ergeben müssen. (Es kommt ungefähr 52.8° heraus.)

Lösung 19.6. Wir ergänzen die Zeichnung um den Punkt D in der Mitte der unteren Seite des Quadrats.



Für den Winkel δ gilt nach Voraussetzung $\tan \delta = 1/2$, weil die Strecke $[BD]$ halb so lang wie $[MD]$ ist. Mit dem Arkustangens erhält man δ . Der Rest ist dann klar, denn offenbar gilt $\gamma + \delta = \pi/4$ und γ ist als Gegenwinkel so groß wie der gesuchte Winkel. (Das Ergebnis liegt bei ungefähr 18.4° .)

Lösung 20.1. Das sollte hoffentlich kein Problem sein:

$$\begin{aligned} \mathbf{a} + \mathbf{b} &= \begin{pmatrix} -2 \\ 5 \end{pmatrix} + \begin{pmatrix} 1 \\ -3 \end{pmatrix} = \begin{pmatrix} -2+1 \\ 5+(-3) \end{pmatrix} = \begin{pmatrix} -1 \\ 2 \end{pmatrix} \\ 3\mathbf{a} &= 3 \cdot \begin{pmatrix} -2 \\ 5 \end{pmatrix} = \begin{pmatrix} 3 \cdot (-2) \\ 3 \cdot 5 \end{pmatrix} = \begin{pmatrix} -6 \\ 15 \end{pmatrix} \\ -\frac{1}{2}\mathbf{b} &= -\frac{1}{2} \cdot \begin{pmatrix} 1 \\ -3 \end{pmatrix} = \begin{pmatrix} (-1/2) \cdot 1 \\ (-1/2) \cdot (-3) \end{pmatrix} = \begin{pmatrix} -1/2 \\ 3/2 \end{pmatrix} \end{aligned}$$

Lösung 20.3. Die Ergebnisse sind $\mathbf{a} + \mathbf{b} = (0, 1)$, $\mathbf{b} - \mathbf{a} = (4, 3)$, $\lambda \cdot \mathbf{a} = (2, 1)$ und $\lambda \cdot \mathbf{b} = (3, 3)$.

Lösung 20.5. Wir wissen bereits, dass $(-1, -3/2)$ auf der Geraden liegt. Mit $\lambda = 1$ (siehe Aufgabe 20.4) erhalten wir einen weiteren Punkt, z. B. $(1/2, -3/4)$. Damit können wir die Steigung berechnen:

$$m = \frac{-\frac{3}{4} - (-\frac{3}{2})}{\frac{1}{2} - (-1)} = \frac{\frac{3}{4}}{\frac{3}{2}} = \frac{1}{2}$$

Da der Punkt $(-1, -3/2)$ auf der Geraden liegt, muss er die Gleichung $y = mx + b$ wahr machen, es muss also

$$-\frac{3}{2} = \frac{1}{2} \cdot (-1) + b$$

gelten. Auflösen liefert $b = -1$. Damit ist $y = x/2 - 1$ die Antwort.

Die Punktmenge g kann man also auch als $\{(x, y) \in \mathbb{R}^2 : y = x/2 - 1\}$ schreiben.

Siehe auch Aufgabe 20.10.

Lösung 20.7. Geraden, die parallel zur y -Achse verlaufen. (Denn welchen Wert sollte die Steigung m dann haben?) Für die Punkt-Richtungs-Form ist das hingegen kein Problem. Als Richtungsvektor könnte man z. B. einfach $(0, 1)$ nehmen.

Lösung 20.8. Mit $\lambda = 2$ erhält man $(2, 0) = (-1, -3/2) + 2 \cdot (3/2, 3/4)$, also liegt $(2, 0)$ auf g . Macht man hingegen den Ansatz

$$\begin{pmatrix} 3 \\ 1 \end{pmatrix} = \begin{pmatrix} -1 \\ -3/2 \end{pmatrix} + \lambda \begin{pmatrix} 3/2 \\ 3/4 \end{pmatrix},$$

so ergibt sich in der ersten Komponente $3 = -1 + \lambda \cdot 3/2$ bzw. $\lambda = 8/3$. Setzt man das in der zweiten Komponente ein, so erhält man $1 = 1/2$, was offensichtlich falsch ist. Daher ist $(3, 1)$ kein Punkt der Geraden.

Lösung 20.9. Aus Aufgabe 20.6 kennen wir andere Punkte der Geraden, beispielsweise $(1/2, -3/4)$ (für $\lambda = 1$). Einen anderen Richtungsvektor kann man z. B. durch $4/3 \cdot (3/2, 3/4) = (2, 1)$ erhalten. Damit wäre

$$g = \left\{ \begin{pmatrix} 1/2 \\ -3/4 \end{pmatrix} + \lambda \begin{pmatrix} 2 \\ 1 \end{pmatrix} : \lambda \in \mathbb{R} \right\}$$

eine weitere Punkt-Richtungs-Form.

Lösung 20.10. Löst man $x = -1 + \lambda \cdot 3/2$ nach λ auf, so ergibt sich $\lambda = 2/3 \cdot (x + 1)$. In der anderen Komponente erhält man durch Einsetzen dieses Wertes:

$$y = -\frac{3}{2} + \frac{2}{3} \cdot (x + 1) \cdot \frac{3}{4} = \frac{x}{2} - 1$$

Das ist natürlich die Geradengleichung aus Aufgabe 20.5.

Lösung 20.11. Wir haben in der ersten Komponente $x = 3 + \lambda \cdot 2$ bzw. $\lambda = 1/2 \cdot (x - 3)$. Einsetzen in die zweite Komponente liefert

$$y = -2 + \frac{1}{2} \cdot (x - 3) \cdot (-1) = -\frac{x}{2} - \frac{1}{2},$$

also $m = b = -1/2$.

Lösung 20.12. Einen Punkt auf der Geraden erhält man ohne Rechnen, wenn man in die Geradengleichung 0 für x einsetzt: $(0, b)$. Dann sollte man sich erinnern, was die Steigung m bedeutet: Wenn man um eine Einheit nach rechts (in x -Richtung) geht, bewegt man sich um m Einheiten nach oben (in y -Richtung). Daher ist $(1, m)$ immer ein Richtungsvektor.

Lösung 20.13. Für die Gerade $y = 12$ kann man als einen Punkt auf der Geraden $(0, 12)$ wählen und für die Richtung Vektor $(1, 0)$. Das ergibt:

$$\begin{pmatrix} 0 \\ 12 \end{pmatrix} + \mathbb{R} \begin{pmatrix} 1 \\ 0 \end{pmatrix}$$

Für $y = -2x + 5$ könnte man in diese Gleichung z. B. $x = 0$ und $x = 1$ einsetzen. Das liefert die beiden Punkte $(0, 5)$ und $(1, 3)$. Damit erhält man:

$$\left\{ \begin{pmatrix} 0 \\ 5 \end{pmatrix} + \lambda \left(\begin{pmatrix} 0 \\ 5 \end{pmatrix} - \begin{pmatrix} 1 \\ 3 \end{pmatrix} \right) : \lambda \in \mathbb{R} \right\} = \begin{pmatrix} 0 \\ 5 \end{pmatrix} + \mathbb{R} \begin{pmatrix} -1 \\ 2 \end{pmatrix}$$

Lösung 20.14. Eine Möglichkeit sieht so aus:

$$\left\{ \begin{pmatrix} 1 \\ 2 \end{pmatrix} + \lambda \left(\begin{pmatrix} -3 \\ 1 \end{pmatrix} - \begin{pmatrix} 1 \\ 2 \end{pmatrix} \right) : \lambda \in \mathbb{R} \right\} = \begin{pmatrix} 1 \\ 2 \end{pmatrix} + \mathbb{R} \begin{pmatrix} -4 \\ -1 \end{pmatrix} \quad (\text{L-3})$$

Man könnte es aber auch so schreiben:

$$\left\{ \lambda \begin{pmatrix} 1 \\ 2 \end{pmatrix} + \mu \begin{pmatrix} -3 \\ 1 \end{pmatrix} : \lambda, \mu \in \mathbb{R} \wedge \lambda + \mu = 1 \right\}$$

Lösung 20.15. Setzt man den Punkt in (L-3) ein, so kann man wie in den vorherigen Aufgaben nach λ auflösen:

$$\begin{pmatrix} 5 \\ 3 \end{pmatrix} = \begin{pmatrix} 1 \\ 2 \end{pmatrix} + (-1) \cdot \begin{pmatrix} -4 \\ -1 \end{pmatrix} \quad (\text{L-4})$$

Daher liegt $(5, 3)$ auf g . Ersetzt man $(-4, -1)$ in (L-4) jedoch durch die Differenz $(-3, 1) - (1, 2)$, so erhält man:

$$\begin{pmatrix} 5 \\ 3 \end{pmatrix} = 2 \cdot \begin{pmatrix} 1 \\ 2 \end{pmatrix} - 1 \cdot \begin{pmatrix} -3 \\ 1 \end{pmatrix}$$

Da weder 2 noch -1 aus dem Intervall $[0, 1]$ sind, liegt der Punkt nicht auf der Verbindungsstrecke von P_1 und P_2 .

Lösung 20.16. Wegen der Darstellung in (20.4) muss man sich lediglich die Faktoren vor den Punkten anschauen. Das sind (in der Reihenfolge der Aufgabenstellung) 1 und 1, 1 und -1 , $2/3$ und $1/3$, $1/3$ und $2/3$ sowie $4/3$ und $-1/3$. In den ersten beiden Fällen ist die Summe der Faktoren *nicht* 1, in den anderen ist sie 1. Daher liegen die drei Punkte rechts auf der Geraden, die beiden links nicht.

Lösung 20.17. In diesem Fall wäre (20.5) ausschlaggebend und man müsste zusätzlich überprüfen, ob die beiden Faktoren im Intervall $[0, 1]$ liegen. Für den letzten Punkt gilt das nicht, für die anderen beiden schon.

Lösung 20.18. Gilt z. B. $\mathbf{b} = \mathbf{0}$, dann ist *jeder* Punkt der Form $\lambda \mathbf{a} + \mu \mathbf{b}$ ein Punkt der Verbindungsgeraden, denn man kann ihn auch als $\lambda \mathbf{a} + (1 - \lambda) \mathbf{b}$ schreiben. Eigentlich handelt es sich einfach um den Punkt $\lambda \mathbf{a}$.

Lösung 20.19. Die Gleichung $(1, a, -3) = \mathbf{p} + \lambda \mathbf{v}$ führt in der ersten Komponente zu $1 = 4 + \lambda \cdot 1$ und damit zu $\lambda = -3$. Setzt man das in die dritte Komponente ein, so ergibt sich $-3 = 6 - 3 \cdot 3$. Nur weil das stimmt, kann ein Punkt der Form $(1, a, -3)$ überhaupt auf g liegen! Mit der zweiten Komponente erhält man nun $a = 5 - 3 \cdot (-2) = 11$.

Lösung 20.20. Man kann das lineare Gleichungssystem z. B. dadurch lösen, dass man die erste Gleichung nach μ auflöst ($\mu = -1 - 2\lambda$) und diesen Wert in die zweite einsetzt:

$$7\lambda - 5(-1 - 2\lambda) = 17\lambda + 5 = 3$$

Daraus folgt $\lambda = -2/17$ und Einsetzen in die zugehörige Geradengleichung liefert den folgenden Schnittpunkt:

$$\mathbf{p} = \begin{pmatrix} 3 \\ -1 \end{pmatrix} - \frac{2}{17} \cdot \begin{pmatrix} 2 \\ 7 \end{pmatrix} = \frac{1}{17} \cdot \begin{pmatrix} 47 \\ -31 \end{pmatrix}$$

Sie können Ihre Lösung überprüfen, indem Sie μ ermitteln ($\mu = -1 - 2 \cdot (-2/17)$) und in die andere Geradengleichung einsetzen. Es muss natürlich derselbe Punkt herauskommen.

Lösung 20.21. Die beiden Geraden können parallel sein. In diesem Fall schneiden sie sich gar nicht, d. h., ihre Schnittmenge ist leer. Oder im Extremfall sind die beiden Geraden nicht nur parallel, sondern sogar identisch, man hat es also eigentlich nur mit einer Geraden zu tun. In diesem Fall besteht die Schnittmenge natürlich aus der ganzen Geraden. Im „normalen“ Fall schneiden die beiden Geraden sich aber in genau einem Punkt wie in der vorherigen Aufgabe.

Machen Sie sich klar, dass es *nicht* möglich ist, dass zwei Geraden mehr als einen Punkt, aber nicht alle gemeinsam haben: wenn zwei Geraden zwei verschiedene Punkte gemeinsam haben, müssen sie bereits identisch sein, da eine Gerade durch zwei Punkte eindeutig festgelegt ist. (Siehe den Anfang von Abschnitt 17.)

Im Raum können auch nur diese drei Fälle (Schnittmenge ist leer, besteht aus einem Punkt oder ist die ganze Gerade) eintreten. Allerdings kann die Schnittmenge auch leer sein, wenn die Geraden *nicht* parallel sind, was in der Ebene nicht möglich ist. Man spricht dann von *windschiefen* Geraden.

windschief

Lösung 20.22. Zum Beispiel so:

$$\begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix} + \mathbb{R} \begin{pmatrix} -4 \\ -1 \\ 5 \end{pmatrix}$$

Lösung 20.23. Sie dürfen nicht parallel sein, denn sonst zeigen beide in dieselbe Richtung und man erhält die Darstellung einer Geraden.

Lösung 20.24. Wie in der Grafik auf Seite 177 zu sehen ist, ist es sinnvoll, beide Richtungsvektoren im selben Punkt „starten“ zu lassen. Man kann also z. B. $\mathbf{p}_1$ als Ortsvektor und $\mathbf{p}_2 - \mathbf{p}_1$ sowie $\mathbf{p}_3 - \mathbf{p}_1$ als Richtungsvektoren wählen. Die resultierende Darstellung wäre dann diese:

$$\begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix} + \mathbb{R} \begin{pmatrix} 2 \\ -3 \\ 4 \end{pmatrix} + \mathbb{R} \begin{pmatrix} 3 \\ -2 \\ -7 \end{pmatrix}$$

Lösung 20.25. Ein Punkt $\mathbf{p}$, der sowohl auf der Geraden als auch auf der Ebene liegt, lässt sich auf zwei Arten darstellen:

$$\mathbf{p} = \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} -4 \\ -1 \\ 5 \end{pmatrix} = \begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix} + \mu \begin{pmatrix} 2 \\ -3 \\ 4 \end{pmatrix} + \nu \begin{pmatrix} 3 \\ -2 \\ -7 \end{pmatrix}$$

Wenn man das komponentenweise aufschreibt, erhält man ein lineares Gleichungssystem mit drei Gleichungen für die drei Unbekannten λ , μ und ν .

$$4\lambda + 2\mu + 3\nu = 0$$

$$-\lambda + 3\mu + 2\nu = 0$$

$$5\lambda - 4\mu + 7\nu = 1$$

Falls Sie sich die Mühe gemacht haben, das auszurechnen: man erhält den Wert $-5/117$ für λ . Als Schnittpunkt ergibt sich $(137/117, 239/117, -25/117)$.

Lösung 20.26. Die Gerade trifft die Ebene entweder in genau einem Punkt oder sie verläuft parallel zu ihr und trifft sie gar nicht oder sie liegt ganz in der Ebene drin.

Zwei Ebenen können identisch sein oder verschieden und parallel (also einen leeren Schnitt haben). Oder sie schneiden sich in einer Geraden. Es ist nicht möglich, dass zwei Ebenen im Raum nur genau einen Punkt gemeinsam haben.²⁸

Lösung 21.2. Das Produkt sieht so aus:

$$\begin{pmatrix} 3 & 7 & 0 \\ -1 & -2 & 9 \\ -1 & 4 & 3 \\ 8 & 0 & -6 \end{pmatrix} \cdot \begin{pmatrix} 5 & 2 & -4 \\ -1 & 3 & 7 \\ 0 & -2 & 2 \end{pmatrix} = \begin{pmatrix} 8 & 27 & 37 \\ -3 & -26 & 8 \\ -9 & 4 & 38 \\ 40 & 28 & -44 \end{pmatrix}$$

Lösung 21.3. $\mathbf{A} \cdot \mathbf{C}$ und $\mathbf{B} \cdot \mathbf{B}$ sind 3×3 -Matrizen, $\mathbf{B} \cdot \mathbf{A}$ ist eine 3×5 -Matrix, $\mathbf{C} \cdot \mathbf{A}$ eine 5×5 -Matrix und $\mathbf{C} \cdot \mathbf{B}$ eine 5×3 -Matrix. Die anderen Produkte kann man alle nicht bilden.

²⁸Anschaulich ist das wohl klar. Formal kann man das damit begründen, dass die Berechnung des Durchschnitts auf ein lineares Gleichungssystem mit drei Gleichungen für vier Unbekannte führt. So ein System nennt man *unterbestimmt*, weil es sozusagen „zu wenige“ Gleichungen hat. Es hat entweder gar keine Lösung oder unendlich viele, aber nicht genau eine.

Lösung 21.4. Die Produkte sind:

$$\mathbf{A} \cdot \mathbf{B} = \begin{pmatrix} 2 & 4 \\ 1 & 2 \end{pmatrix} \quad \mathbf{B} \cdot \mathbf{A} = \begin{pmatrix} 4 & 2 \\ 0 & 0 \end{pmatrix}$$

Lösung 21.5. Nein. Sie haben hoffentlich sofort bemerkt, dass Aufgabe 21.4 diese Frage schon beantwortet.

Lösung 21.6. Das Ergebnis:

$$\begin{pmatrix} 5 & 5 & 2 \\ 0 & 1 & 2 \\ 5 & 3 & 6 \end{pmatrix}$$

Lösung 21.7. Die Antwort folgt im Text.

Lösung 21.8. Beim Multiplizieren stellt man fest, dass die beiden Nullen dafür sorgen, dass $(a^2, 0, 0)$ die erste Zeile von $\mathbf{M}^2$ ist. Weil wir nun wieder zwei Nullen an derselben Stelle haben, wird es offenbar so weitergehen. Die gesuchte Zahl ist daher a^4 . (Hoffentlich haben Sie nicht nur die Lösung gelesen, sondern die Multiplikationen zumindest für die jeweils erste Zeile wirklich durchgeführt!)

Lösung 21.9. Es klappt, wenn man bei $\mathbf{E}_2$ eine der beiden Einsen weglässt.

Lösung 21.10. Es ergibt sich das folgende Produkt:

$$\begin{pmatrix} 4 & 5 & 0 \\ -2 & -2 & 8 \\ -1 & 4 & 3 \\ 7 & 0 & -6 \end{pmatrix} \cdot \begin{pmatrix} 2 \\ 3 \\ -2 \end{pmatrix} = \begin{pmatrix} 23 \\ -26 \\ 4 \\ 26 \end{pmatrix}$$

Lösung 21.11. Das ist eigentlich eher eine Aufgabe für die Grundschule. Offenbar muss $a = 1$, $b = 2$ und $c = 7$ gelten.

Lösung 21.12. Hier muss man ein bisschen mehr nachdenken als bei der Aufgabe davor, aber das Prinzip ist dasselbe. Es ist ein bisschen wie **Sudoku**. Zunächst sieht man, dass $a = 2$ gelten muss. Damit folgt $b = 2a = 4$ und schließlich $d = a - 2b = -6$.

Lösung 21.15. Man bekommt wieder einen transponierten Vektor. Wegen (21.2) ist das Ergebnis auch der transponierte Vektor der „normalen“ Multiplikation:

$$\mathbf{v}^T \cdot \mathbf{A} = (\mathbf{A}^T \cdot \mathbf{v})^T$$

Lösung 21.16. Man kann dieselbe Lösung wie in Aufgabe 21.9 nehmen.

Lösung 22.1. Die erste Gleichung hat die zwei Lösungen -5 und 5 . In der zweiten kann man für y die Werte -4 und 4 einsetzen. Für das Paar (x, y) gibt es also vier Lösungen:

$$(5, 4) \quad (5, -4) \quad (-5, 4) \quad (-5, -4)$$

Lösung 22.2. Die Matrix sieht so aus:

$$\left(\begin{array}{ccc|c} 6 & -1 & -1 & 4 \\ 1 & 1 & 10 & -6 \\ 2 & -1 & 1 & -2 \end{array} \right)$$

Lösung 22.3. $\mathbf{M}$ war die Koeffizientenmatrix aus Aufgabe 22.2. Der Sinn dieser Aufgabe ist, zu sehen, dass das LGS aus der besagten Aufgabe sich auch in der Form $\mathbf{M} \cdot \mathbf{x} = \mathbf{b}$ mit $\mathbf{b} = (4, -6, -2)$ schreiben lässt, wobei die drei Komponenten die drei Gleichungen liefern. Die Koeffizientenmatrix ist also nicht nur eine platzsparende Schreibweise, sondern man kann ein LGS auch als Matrizengleichung interpretieren, bei der die Unbekannte ein Vektor ist.

Lösung 22.4. Die dritte Gleichung liefert $z = -3$. Einsetzen in die zweite Gleichung ergibt $2y - 3 = 7$ bzw. $y = 5$. Mit diesen beiden Werten ergibt sich in der ersten Gleichung $-x + 15 - 3 = 13$, also $x = -1$. Die einzige Lösung ist somit $(-1, 5, -3)$, die Lösungsmenge also $\{(-1, 5, -3)\}$.

Lösung 22.5. Das könnte so aussehen:

$$\left(\begin{array}{ccc} 0 & \textcircled{1} & 1 \\ 0 & \textcircled{1} & 1 \\ 0 & 0 & \textcircled{1} \end{array} \right)$$

Dies ist wegen der blauen Nullen eine obere Dreiecksmatrix, aber der Leitkoeffizient der zweiten Zeile steht nicht weiter rechts als der der ersten Zeile, also hat die Matrix nicht Zeilenstufenform.

Lösung 22.6. Zunächst sei darauf hingewiesen, dass es beim Gaußverfahren viele Wege gibt, die nach Rom führen. Wenn Sie die Aufgabe anders gelöst haben als ich, ist das kein Problem. Auch die Matrix in Zeilenstufenform am Ende kann eine andere sein. Allerdings muss die Lösungsmenge dieselbe sein, die auch ich hier erhalte, sonst hat einer von uns beiden einen Fehler gemacht!

Wir beginnen mit dieser Matrix.

$$\left(\begin{array}{ccc|c} 6 & -1 & -1 & 4 \\ 1 & 1 & 10 & -6 \\ 2 & -1 & 1 & -2 \end{array} \right)$$

Wir haben eigentlich schon ein Pivotelement an der richtigen Stelle. Wenn man ohne elektronische Hilfsmittel rechnet, ist es aber oft hilfreich, ein Pivotelement zu wählen, mit dem sich leicht rechnen lässt. Das geht besonders gut mit den Zahlen 1 und -1 oder ersatzweise mit ganzen Zahlen, die nicht weit von der Null entfernt sind. Ich vertausche daher zunächst die ersten beiden Zeilen.

$$\left(\begin{array}{ccc|c} \textcircled{1} & 1 & 10 & -6 \\ 6 & -1 & -1 & 4 \\ 2 & -1 & 1 & -2 \end{array} \right)$$

Nun subtrahiere ich das Sechsfache der ersten von der zweiten Zeile.

$$\left(\begin{array}{ccc|c} \textcircled{1} & 1 & 10 & -6 \\ 0 & -7 & -61 & 40 \\ 2 & -1 & 1 & -2 \end{array}\right)$$

Anschließend wird das Doppelte der ersten von der dritten Zeile abgezogen.²⁹

$$\left(\begin{array}{ccc|c} 1 & 1 & 10 & -6 \\ 0 & -7 & -61 & 40 \\ 0 & -3 & -19 & 10 \end{array}\right)$$

Nun hat man keine „schöne“ Zahl mehr, die man als Pivotelement wählen kann. Nimmt man z. B. -7 , so muss man die zweite Zeile mit $3/7$ multiplizieren und von der dritten subtrahieren, um eine Null zu erzeugen. Dabei kommen unangenehme Brüche ins Spiel.

Man kann das aber vermeiden (auf Kosten größerer Zahlen), indem man vorher „prophylaktisch“ die dritte Zeile mit 7 multipliziert.

$$\left(\begin{array}{ccc|c} 1 & 1 & 10 & -6 \\ 0 & \textcircled{-7} & -61 & 40 \\ 0 & -21 & -133 & 70 \end{array}\right)$$

Jetzt kann man das Dreifache der zweiten von der dritten Zeile subtrahieren, ohne dass irgendwo Brüche auftauchen.

$$\left(\begin{array}{ccc|c} 1 & 1 & 10 & -6 \\ 0 & -7 & -61 & 40 \\ 0 & 0 & 50 & -50 \end{array}\right)$$

Wir haben eine Zeilenstufenform erreicht und können durch „Rückwärtseinsetzen“ wie in Aufgabe 22.4 die Lösung finden. Die Lösungsmenge ist $\{(1, 3, -1)\}$.

Lösung 22.7. Im ersten Schritt ist die Eins links oben unser Pivotelement und wir erledigen die zweite und die dritte Zeile in einem Schritt.

$$\left(\begin{array}{ccc|c} \textcircled{1} & 0 & 1 & 5 \\ 0 & 1 & 1 & -2 \\ 0 & 1 & 1 & -3 \end{array}\right)$$

Nun haben wir netterweise wieder eine Eins als Pivotelement und können die Null unter ihr auch ganz einfach erzeugen.

$$\left(\begin{array}{ccc|c} 1 & 0 & 1 & 5 \\ 0 & \textcircled{1} & 1 & -2 \\ 0 & 0 & 0 & -1 \end{array}\right)$$

²⁹Die beiden letzten Schritte hätte man auch in einem zusammenfassen können. Das bietet sich insbesondere dann an, wenn solche Berechnung von Hand durchgeführt werden und Schreibarbeit gespart werden soll. Man muss dabei aber aufpassen, dass man die elementaren Zeilenumformungen immer noch *nacheinander* ausführt und lediglich beim Aufschreiben Schritte spart, sonst macht man evtl. Fehler.

Lösung 22.8. Um uns die Arbeit zu erleichtern, vertauschen wir zunächst die erste mit der dritten Zeile. Dann werden in den Zeilen darunter Nullen erzeugt.

$$\left(\begin{array}{ccc|c} \textcircled{-1} & -1 & 9 & -12 \\ 0 & -1 & 16 & -18 \\ 0 & -2 & 32 & -36 \end{array} \right)$$

Wir fahren fort mit der zweiten Zeile und sind in einem Schritt schon fertig:

$$\left(\begin{array}{ccc|c} -1 & -1 & 9 & -12 \\ 0 & \textcircled{-1} & 16 & -18 \\ 0 & 0 & 0 & 0 \end{array} \right)$$

Lösung 22.9. Das ist die **Punkt-Richtungs-Form** einer Geraden, und zwar

$$(-6, 18, 0) + \mathbb{R}(-7, 16, 1).$$

Wie wir noch sehen werden, ist das kein Zufall, weil man LGS geometrisch interpretieren kann.

Lösung 22.10. Ja, das geht. In unserem ersten Beispiel hatten wir es etwa mit dieser Matrix zu tun, in deren erster Zeile der Leitkoeffizient zu weit rechts stand:

$$\left(\begin{array}{cccc} 0 & 0 & -1 & 5 \\ 2 & -2 & 1 & 3 \\ 4 & -4 & 3 & 8 \end{array} \right)$$

Statt die erste Zeile mit einer anderen zu vertauschen, können wir aber z. B. einfach die zweite zur ersten addieren:

$$\left(\begin{array}{cccc} \textcircled{2} & -2 & 0 & 8 \\ 2 & -2 & 1 & 3 \\ 4 & -4 & 3 & 8 \end{array} \right)$$

Nun ist der Leitkoeffizient in der ersten Zeile an der richtigen Stelle und wir können wie üblich fortfahren. Man kann also beim Gaußverfahren ausschließlich mit elementaren Zeilenumformungen vom Typ (III) auskommen.

Lösung 22.11. Die Lösungsmenge ist $\{(4, -3, 2, -1)\}$.

Sie können (und sollten) sich natürlich selbst Aufgaben stellen. Ob Ihre Lösung richtig ist, können Sie in der Regel einfach durch Einsetzen überprüfen. Man kann aber z. B. auch in **Wolfram Alpha** so etwas wie

$$\text{solve } x-z=5, \ 2x+y+3z=7, \ 3x+y+4z=12$$

eingeben und sich die Antwort anschauen.

Lösung 22.12. Die Lösungsmenge ist $\{(4, 4, 3)\}$. Als Komponenten dürfen natürlich nur Elemente von $\mathbb{Z}_7 = \{0, 1, \dots, 6\}$ vorkommen.

Lösung 22.13. Die Lösungsmenge ist $\{(1, 0, 1)\}$. Auch hier gehen Sie wie bisher vor. Sie sollten lediglich beachten, dass es keine gute Idee ist, a als Pivotelement zu wählen, weil der Fall $a = 0$ nicht ausgeschlossen ist.

In diesem Fall hängt die Lösung nicht von a ab, aber das kann durchaus passieren.

Lösung 22.14. Ich zeige die Lösungen mithilfe von Beispielen, wobei $\mathbf{A}$ jeweils eine $4 \times n$ -Matrix sein soll. Für den Typ (I) multiplizieren wir die dritte Zeile von $\mathbf{A}$ mit 7:

$$\text{3. Zeile} \rightarrow \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 7 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \cdot \mathbf{A}$$

Und für Typ (II) vertauschen wir die erste Zeile mit der dritten:

$$\begin{array}{l} \text{1. Zeile} \rightarrow \begin{pmatrix} 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \cdot \mathbf{A} \\ \text{3. Zeile} \rightarrow \end{array}$$

$\uparrow$
1. Spalte

$\uparrow$
3. Spalte

Das demonstriert hoffentlich, wie man es im allgemeinen Fall machen würde.

Lösung 23.1. $f(\mathbf{p})$ und $f(\mathbf{q})$ könnten derselbe Punkt sein. Dann würde die gesamte Gerade durch $\mathbf{p}$ und $\mathbf{q}$ auf einen Punkt abgebildet werden.

Lösung 23.2. Das sieht so aus:

$$f\left(\alpha \cdot \begin{pmatrix} v_1 \\ v_2 \end{pmatrix}\right) = f\left(\begin{pmatrix} \alpha v_1 \\ \alpha v_2 \end{pmatrix}\right) = \begin{pmatrix} 2\alpha v_1 + 3\alpha v_2 \\ 3\alpha v_1 - \alpha v_2 \end{pmatrix} = \alpha \cdot \begin{pmatrix} 2v_1 + 3v_2 \\ 3v_1 - v_2 \end{pmatrix} = \alpha \cdot f\left(\begin{pmatrix} v_1 \\ v_2 \end{pmatrix}\right)$$

Lösung 23.3. Das sind die Vektoren $(2, 3)$ und $(3, -1)$.

Lösung 23.4. Durch Einsetzen in die Definition der Funktion erhält man die Werte $f(\mathbf{e}_1) = (5, 4)$, $f(\mathbf{e}_2) = (-2, 0)$ und $f(\mathbf{e}_3) = (0, -1)$. Nach Lemma 23.2 sind das die Spalten der Matrix, so dass wir

$$\mathbf{A} = \begin{pmatrix} 5 & -2 & 0 \\ 4 & 0 & -1 \end{pmatrix}$$

erhalten.

Lösung 23.5. Die Spalten der Matrix sind bereits angegeben. Der gesuchte Funktionswert ist

$$f((3, 2)) = \begin{pmatrix} 2 & 2 \\ 1 & -4 \end{pmatrix} \cdot \begin{pmatrix} 3 \\ 2 \end{pmatrix} = \begin{pmatrix} 10 \\ -5 \end{pmatrix}.$$

Lösung 23.6. Wegen der Linearität von f gilt

$$f(\mathbf{e}_2) = f\left(\frac{1}{2} \cdot 2\mathbf{e}_2\right) = \frac{1}{2} \cdot f(2\mathbf{e}_2) = \frac{1}{2} \cdot \begin{pmatrix} 2 \\ -4 \end{pmatrix} = \begin{pmatrix} 1 \\ -2 \end{pmatrix}.$$

Damit hat man die zweite Spalte der Matrix und kann $f((3, 2)) = (8, -1)$ berechnen.

Lösung 23.7. Sie wollen α und β mit $(3, 1) = \alpha \cdot (2, 1) + \beta \cdot (0, 2)$ finden. Das ergibt das lineare Gleichungssystem, das aus den Gleichungen $3 = 2\alpha$ und $1 = \alpha + 2\beta$ besteht. Die erste Gleichung liefert $\alpha = 3/2$ und die zweite $\beta = (1 - \alpha)/2 = -1/4$. Es folgt $f((3, 1)) = 3/2 \cdot f((2, 1)) - 1/4 f((0, 2)) = 3/2 \cdot (3, 2) - 1/4 \cdot (1, -5) = 17/4 \cdot (1, 1)$.

Lösung 23.8. Das Prinzip ist dasselbe wie bei den vorherigen Aufgaben: man muss sich das „zusammenbasteln“. Bei komplizierteren Problemen stellt man ein ggf. ein **lineares Gleichungssystem** wie $\mathbf{e}_3 = \alpha\mathbf{e}_1 + \beta\mathbf{e}_2 + \gamma\mathbf{v}_3$ auf, aber in diesem Fall sollte man das direkt sehen können.

$$\begin{aligned} f(\mathbf{e}_3) &= f(\mathbf{v}_3 - 2\mathbf{e}_1 + \mathbf{e}_2) = f(\mathbf{v}_3) - 2f(\mathbf{e}_1) + f(\mathbf{e}_2) = \mathbf{v}_4 - 2\mathbf{v}_1 + \mathbf{v}_2 \\ f(\mathbf{w}) &= f(3\mathbf{e}_1 + 2\mathbf{e}_2 + \mathbf{e}_3) = 3f(\mathbf{e}_1) + 2f(\mathbf{e}_2) + f(\mathbf{e}_3) = \\ &= 3\mathbf{v}_1 + 2\mathbf{v}_2 + (\mathbf{v}_4 - 2\mathbf{v}_1 + \mathbf{v}_2) = \mathbf{v}_1 + 3\mathbf{v}_2 + \mathbf{v}_4 \end{aligned}$$

Insgesamt ergibt sich $f(\mathbf{w}) = (-9, 16, 27)$.

Lösung 23.9. Abgesehen davon, dass man beim Lesen der Aufgabe nicht übersehen sollte, dass hier in $\mathbb{Z}_5$ gerechnet wird, läuft alles wie in den Aufgaben vorher.

$$\begin{aligned} f(\mathbf{e}_1) &= f(2 \cdot 3\mathbf{e}_1) = 2f(3\mathbf{e}_1) = 2 \cdot (2, 4, 3) = (4, 3, 1) \\ f(\mathbf{e}_2) &= f(3 \cdot 2\mathbf{e}_2) = 3f(2\mathbf{e}_2) = 3 \cdot (1, 0, 3) = (3, 0, 4) \\ f((4, 1)) &= f(4\mathbf{e}_1 + \mathbf{e}_2) = 4f(\mathbf{e}_1) + f(\mathbf{e}_2) = 4 \cdot (4, 3, 1) + (3, 0, 4) = (4, 2, 3) \end{aligned}$$

Lösung 23.10. Die Matrizen zu den Abbildungen sind diese:

$$\mathbf{A}_f = \begin{pmatrix} 2 & -3 & 0 \\ -1 & 0 & 1 \end{pmatrix} \quad \text{und} \quad \mathbf{A}_g = \begin{pmatrix} 1 & 0 & 2 \\ 0 & 5 & -1 \\ 1 & -1 & 1 \end{pmatrix}$$

Die Matrix zu $f \circ g$ sieht folglich so aus:

$$\mathbf{A}_f \cdot \mathbf{A}_g = \begin{pmatrix} 2 & -15 & 7 \\ 0 & -1 & -1 \end{pmatrix}$$

$f \circ g$ ist eine Abbildung von $\mathbb{R}^3$ nach $\mathbb{R}^2$. $g \circ f$ kann man hingegen nicht bilden, weil der Wertebereich von f die Ebene $\mathbb{R}^2$ ist, der Definitionsbereich von g aber der Raum $\mathbb{R}^3$. Man erkennt das auch daran, dass man das Produkt $\mathbf{A}_g \cdot \mathbf{A}_f$ nicht bilden kann.

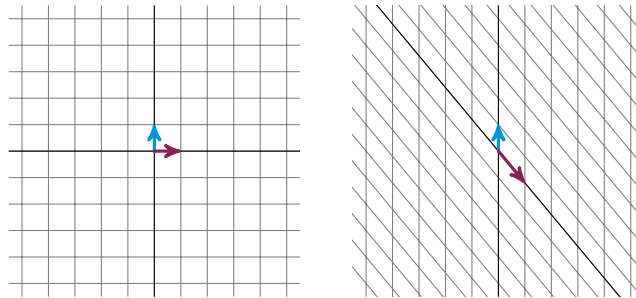
Lösung 23.11. f bildet $(1, 0) = \frac{1}{2} \cdot (2, 0)$ auf $\frac{1}{2} \cdot (2, 4) = (1, 2)$ ab. Damit hat man die Matrix für f und erhält als Matrix für $g \circ f$

$$\begin{pmatrix} 3 & -1 \\ 2 & 0 \end{pmatrix} \cdot \begin{pmatrix} 1 & 2 \\ 2 & 2 \end{pmatrix} = \begin{pmatrix} 1 & 4 \\ 2 & 4 \end{pmatrix}.$$

Lösung 24.1. $\begin{pmatrix} 3.7 & 0 \\ 0 & 3.7 \end{pmatrix}$

Lösung 24.3. $\mathbf{M}$ muss die $n \times n$ -Diagonalmatrix sein, in deren Hauptdiagonale überall λ steht, also $\mathbf{M} = \lambda \mathbf{E}_n$.

Lösung 24.4. In diesem Fall wird längs der vertikalen Achse geschert. Und nach „unten“, weil -1.2 negativ ist.



Lösung 24.5. Bezeichnet man mit $\mathbf{D}_\alpha$ die Matrix für die Drehung um den Winkel α , dann sollte Ihre Rechnung so ausgesehen haben:

$$\begin{aligned} \mathbf{D}_\alpha \cdot \mathbf{D}_\beta &= \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \cdot \begin{pmatrix} \cos \beta & -\sin \beta \\ \sin \beta & \cos \beta \end{pmatrix} \\ &= \begin{pmatrix} \cos \alpha \cos \beta - \sin \alpha \sin \beta & -\sin \alpha \cos \beta - \sin \beta \cos \alpha \\ \sin \alpha \cos \beta + \sin \beta \cos \alpha & \cos \alpha \cos \beta - \sin \alpha \sin \beta \end{pmatrix} \\ &= \begin{pmatrix} \cos(\alpha + \beta) & -\sin(\alpha + \beta) \\ \sin(\alpha + \beta) & \cos(\alpha + \beta) \end{pmatrix} = \mathbf{D}_{\alpha+\beta} \end{aligned}$$

Lösung 24.6. Grundsätzlich sollten Sie sich bei solchen Fragen zuerst eine Skizze machen und wie in den entsprechenden Abbildungen im Skript die Bilder der Einheitsvektoren einzeichnen, die man direkt der Matrix entnehmen kann. Man sieht dann hoffentlich sofort: Die erste Matrix beschreibt eine Drehung um 90° . Die zweite ist die Spiegelungsmatrix für den Winkel 45° , d. h., es wird an der Hauptdiagonalen gespiegelt. Das kann man nun zur Sicherheit noch überprüfen, indem man die entsprechenden Matrizen ausrechnet.

Lösung 24.7. In beiden Fällen kann man sich sofort (ggf. mithilfe einer Skizze) überlegen, auf welche Vektoren die beiden kanonischen Basisvektoren abgebildet werden müssen. Das müssen die Spalten der jeweiligen Matrix sein. Das Ergebnis sieht also so aus:

$$\begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix} \quad \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$$

Lösung 24.8. Die Matrix für eine Drehung um 180° kennen wir schon aus Aufgabe 24.7:

$$\mathbf{D} = \begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix}$$

Durch Nachrechnen überprüft man nun, dass tatsächlich $\mathbf{D} \cdot \mathbf{D} = \mathbf{E}_2$ stimmt.

Die Matrix für das Spiegeln an der y -Achse sieht so aus:

$$\mathbf{S} = \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix}$$

Auch hier ergibt sich sofort $\mathbf{S} \cdot \mathbf{S} = \mathbf{E}_2$.

Lösung 24.9. Da der Mittelpunkt des Quadrates der Koordinatenursprung ist, haben wir es offenbar mit einer Drehung des Quadrates zu tun, das die Ecken $(\pm 1, \pm 1)$ hat. Vor der Drehung wäre der Winkel der Linie vom Ursprung zur oberen rechten Ecke zur horizontalen Achse 45° gewesen. Es wurde somit um $57 - 45 = 12$ Grad gedreht. Die gesuchten Koordinaten erhält man, in dem man den Punkt $(1, -1)$ mit der Drehmatrix für 12° multipliziert. Alternativ kann man auch direkt argumentieren, dass der Punkt P die Polarkoordinaten $\sqrt{2}$ und $-45 + 12$ Grad haben muss. (Warum?) Es ergibt sich jedenfalls gerundet $(1.19, -0.77)$.

(Beachten Sie, dass man die ungefähre Lage von P der Skizze entnehmen kann. Sollten Sie beispielsweise eine negative x -Koordinate herausbekommen haben, dann sollte schon deswegen klar sein, dass Sie einen Fehler gemacht haben müssen.)

Lösung 24.10. Sei $\mathbf{D}_\alpha$ wie in der Lösung von Aufgabe 24.5 die entsprechende Drehmatrix. Der Winkel, den die Gerade, auf die projiziert werden soll, mit der x -Achse bildet, ist $\alpha = \arctan(9/7)$. Dreht man um $-\alpha$, so wird die Gerade also auf die horizontale Achse gedreht, und die Projektion wird dann durch die Matrix (24.2) dargestellt. Insgesamt erhält man die gesuchte Matrix damit durch

$$\mathbf{D}_\alpha \cdot \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \cdot \mathbf{D}_{-\alpha} = \frac{1}{130} \cdot \begin{pmatrix} 49 & 63 \\ 63 & 81 \end{pmatrix} \approx \begin{pmatrix} 0.38 & 0.48 \\ 0.48 & 0.62 \end{pmatrix}.$$

(Die Matrix in der Mitte erhält man, wenn man exakt rechnet, z. B. mit einem Computeralgebrasystem, die Matrix rechts mit einem Taschenrechner. Die Zahlen sind aber nicht so wichtig. Entscheidend ist, dass Sie wissen, was man in den Taschenrechner eintippen müsste, dass Sie also auf die Matrix links kommen.)

Lösung 24.11. Der Winkel der Geraden aus der Aufgabenstellung mit der horizontalen Achse ist $\alpha = \arctan(1/3)$. Die Matrix für die Spiegelung ist daher

$$\begin{pmatrix} \cos 2\alpha & \sin 2\alpha \\ \sin 2\alpha & -\cos 2\alpha \end{pmatrix} = \frac{1}{5} \cdot \begin{pmatrix} 4 & 3 \\ 3 & -4 \end{pmatrix}.$$

Lösung 24.12. Die Spiegelungsmatrizen sehen so aus:

$$\mathbf{S}_x = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad \mathbf{S}_y = \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix}$$

Es ergibt sich:

$$\mathbf{S}_x \mathbf{S}_y = \begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix} = \mathbf{S}_y \mathbf{S}_x$$

Es ist also egal, welche von den beiden Spiegelungen man zuerst durchführt.

Lösung 25.1. Ihre Rechnung sollte im Wesentlichen so ausgesehen haben:

$$\begin{bmatrix} a & b & 0 \\ c & d & 0 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} x \\ y \\ 1 \end{bmatrix} = \begin{bmatrix} ax + by \\ cx + dy \\ 1 \end{bmatrix}$$

$$\begin{bmatrix} 1 & 0 & a \\ 0 & 1 & b \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} x \\ y \\ 1 \end{bmatrix} = \begin{bmatrix} x + a \\ y + b \\ 1 \end{bmatrix}$$

Wenn man die dritte Komponente des Ergebnisses (die offensichtlich immer die Zahl 1 ist) ignoriert, kommt genau das heraus, was man erwarten würde.

Lösung 25.2. Die Antwort wird bereits durch die Verwendung des Grautons angedeutet: Solange wir nur die beiden Typen (25.1) und (25.2) einsetzen, werden Produkte von solchen Matrizen immer dieselbe letzte Zeile haben. Daher kann man diese auch weglassen. Und genau das machen manche Softwarebibliotheken auch.

Lösung 25.3. Die Matrizen für f und h kann man ohne großes Nachdenken hinschreiben (siehe Aufgabe 24.7) und in (25.1) einsetzen. Ebenso sollte die Matrix für g nach dem Muster (25.2) hoffentlich ein Kinderspiel sein. Insgesamt erhält man also

$$\begin{bmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} 1 & 0 & 3 \\ 0 & 1 & 2 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} 0 & -1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 1 \end{bmatrix} = \begin{bmatrix} 0 & 1 & -3 \\ 1 & 0 & 2 \\ 0 & 0 & 1 \end{bmatrix}$$

Lösung 25.4. g ist offenbar eine Parallele zur y -Achse und die Matrix für die Spiegelung an der y -Achse kann man wie in Aufgabe 24.7 ohne Rechnen angeben. Analog zum vorgerechneten Beispiel verschieben wir erst um $(-2, 0)$, um g mit der y -Achse in Deckung zu bringen. Dann wird gespiegelt und anschließend wird die Translation zurückgenommen. Die Matrix sieht also folgendermaßen aus:

$$\begin{bmatrix} 1 & 0 & 2 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} 1 & 0 & -2 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} = \begin{bmatrix} -1 & 0 & 4 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

Lösung 26.1. Nach unseren Vorüberlegungen müssen wir eine Matrix finden, deren zugehörige lineare Abbildung nicht injektiv ist, weil sie (mindestens) eine Gerade auf einen Punkt abbildet. Die einfachste Antwort ist die Nullmatrix, denn sie gehört zu der Abbildung, die sämtliche Punkte der Ebene auf den Nullpunkt abbildet.

Man könnte aber auch die Projektion aus Abschnitt 24 als Beispiel nehmen:

$$A = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$$

Geraden, die parallel zur y -Achse sind, werden durch diese Projektion auf einen Punkt abgebildet, daher ist sie nicht injektiv und die Matrix kann nicht regulär sein.

Man kann sogar „zu Fuß“ leicht ausrechnen, dass **A** nicht invertierbar ist. Dazu wählen wir uns eine beliebige 2×2 -Matrix und multiplizieren diese mit **A**:

$$\begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \cdot \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{pmatrix} = \begin{pmatrix} b_{11} & b_{12} \\ 0 & 0 \end{pmatrix}$$

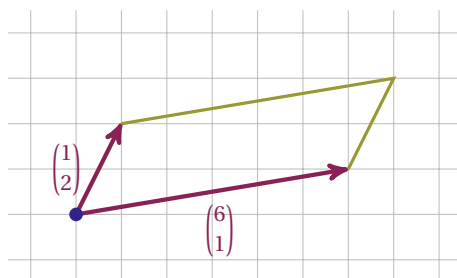
Völlig unabhängig davon, welche Werte in der Matrix stehen, an der rot markierten Stelle wird immer eine Null auftauchen. Daher kann rechts vom Gleichheitszeichen nie die Einheitsmatrix stehen.

Lösung 27.1. Das große Rechteck hat die Fläche $(a+b)(c+d)$. Die beiden kleinen roten Rechtecke haben jeweils die Fläche bc . Die beiden grünen Dreiecke haben zusammen die Fläche bd und für die beiden blauen ergibt sich ac . Insgesamt kommt man auf

$$\begin{aligned} (a+b)(c+d) - 2bc - bd - ac &= ac + bc + ad + bd - 2bc - bd - ac \\ &= ad - bc. \end{aligned}$$

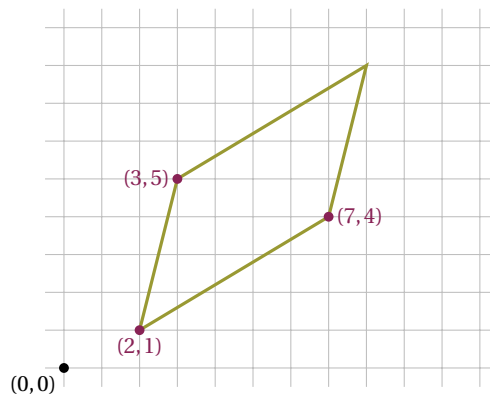
Lösung 27.2. Es ergibt sich $5 \cdot (-3) - 2 \cdot (-8) = 1$ für die erste Determinante und $42 \cdot 0 - 1/3 \cdot 6 = -2$ für die zweite.

Lösung 27.3. Legt man den Ursprung eines gedachten Koordinatensystems auf den blauen Punkt, so kann man die Koordinaten der beiden (roten) Vektoren ablesen, die das Parallelogramm aufspannen.



Als Fläche ergibt sich somit $6 \cdot 2 - 1 \cdot 1 = 11$.

Lösung 27.4. Grafisch könnte es so aussehen:



Man geht nun wie in der vorherigen Aufgabe vor und verschiebt den Ursprung auf den Punkt $(2, 1)$. Als aufspannende Vektoren des Parallelogramms erhält man:

$$\begin{pmatrix} 7 \\ 4 \end{pmatrix} - \begin{pmatrix} 2 \\ 1 \end{pmatrix} = \begin{pmatrix} 5 \\ 3 \end{pmatrix} \quad \text{und} \quad \begin{pmatrix} 3 \\ 5 \end{pmatrix} - \begin{pmatrix} 2 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ 4 \end{pmatrix}$$

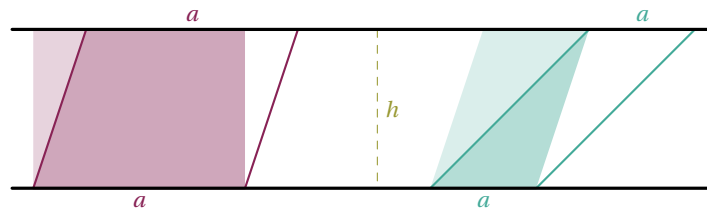
Daher ist die Fläche $5 \cdot 4 - 1 \cdot 3 = 17$. (Beachten Sie, dass wir zur Berechnung der Fläche die Koordinaten der vierten Ecke gar nicht zu ermitteln brauchten.)

Lösung 27.5. Nach Aufgabenstellung soll $\mathbf{a}$ die Form $(a, 0)$ haben. Die Fläche des Parallelogramms ist dann nach (27.1) $2a$. Also muss $a = 5/2$ gelten.

Lösung 27.6. Die beiden Spaltenvektoren sind parallel, also gibt es nur ein „entartetes“ Parallelogramm, das keinen Flächeninhalt hat.

Lösung 27.7. Dieselbe, denn offenbar gilt $ad - bc = ad - cb$.

Lösung 27.8. Man kann das beispielsweise wie in der folgenden Abbildung machen, in der die Fläche sich jeweils als ah ergibt.



Links sieht man, dass das Parallelogramm dieselbe Fläche wie das Rechteck hat, weil man auf der einen Seite ein Dreieck (weiß) entfernt, das man auf der anderen (helles Rot) wieder hinzufügt. Das rechte Parallelogramm ist zu „schief“, um denselben „Trick“ anzuwenden. Man kann es jedoch durch Wegnehmen und Hinzufügen eines Dreiecks „weniger schief“ machen, ohne seine Fläche zu ändern. (Bei besonders „schiefen“ Parallelogrammen muss man das evtl. mehrfach machen.)

Lösung 27.9. Die Lösungen sind $\det(\mathbf{A}) = -260$ und $\det(\mathbf{B}) = -112$, falls ich mich nicht vertippt habe.

Lösung 27.10. Die erste Determinante ist:

$$1 \cdot 7 \cdot (-4) + 3 \cdot 0 \cdot 2 + 5 \cdot 8 \cdot 1 - 5 \cdot 7 \cdot 2 - 3 \cdot 8 \cdot (-4) - 1 \cdot 0 \cdot 1 = 38$$

Im zweiten Fall ergibt sich null. (Die drei Vektoren liegen nämlich alle in einer Ebene, daher hat der Spat kein Volumen.)

Lösung 27.11. Das Ergebnis ist 1. Wie wir in Abschnitt 9 gelernt haben, kann man einerseits erst in $\mathbb{R}$ rechnen und aus dem Ergebnis 100 dann $100 \bmod 11 = 1$ machen oder man kann zwischendurch bereits einzelne Ergebnisse „kleiner machen“, um den Rechenaufwand zu verringern.

Lösung 27.12. $\mathbf{A}$ beschreibt eine Drehung. Für die Determinante ergibt sich nach Pythagoras $\cos^2(\pi/5) + \sin^2(\pi/5) = 1$. Das entspricht der Tatsache, dass Drehungen weder Flächen (bzw. Volumen) noch Orientierungen ändern.

B beschreibt eine Spiegelung und als Determinante ergibt sich -1 . Spiegelungen ändern Flächeninhalte nicht, aber sie ändern Orientierungen.

C ist die Matrix einer Scherung längs der x -Achse. Die Determinante ist 1, denn so eine Scherung verändert Flächen nicht. (Falls Ihnen das nicht ohnehin klar ist, zeichnen Sie sich das Parallelogramm, in das das von den kanonischen Einheitsvektoren aufgespannte Quadrat durch die Scherung überführt wird.)

D schließlich beschreibt eine Streckung in x -Richtung um den Faktor 3. Dadurch verdreifachen sich alle Flächen und entsprechend hat die Determinante von **D** auch den Wert 3.

Die Determinanten von **C** und **D** sind natürlich beide positiv, weil sich keine Orientierungen ändern.

Lösung 27.13. Eine Einheitsmatrix hat natürlich immer die Determinante 1. Es handelt sich um eine Dreiecksmatrix, also muss man einfach das Produkt der Diagonaleinträge bilden; und das sind alles Einsen.

Das ist aber auch geometrisch klar, weil eine Einheitsmatrix ja eine Abbildung beschreibt, die nichts ändert. Daher kann der zugehörige Vergrößerungsfaktor nur eins sein.

Lösung 27.14. Man könnte z. B. $-\mathbf{E}_2$ nehmen. Die Determinante dieser Matrix ist 1, denn weil beide Spalten der Matrix $\mathbf{E}_2$ mit -1 multipliziert werden, wird die Determinante von $\mathbf{E}_2$ zweimal mit -1 multipliziert, wodurch sich insgesamt der Faktor 1 ergibt. (Siehe dazu auch Aufgabe 27.15 mit $\lambda = -1$.)

Lösung 27.15. Es gilt $\det(\lambda \cdot \mathbf{A}) = \lambda^n \cdot \det(\mathbf{A})$. Das liegt daran, dass man *jede* der n Zeilen (oder Spalten) von $\mathbf{A}$ mit λ multiplizieren muss, um aus $\mathbf{A}$ die Matrix $\lambda \cdot \mathbf{A}$ zu machen.

Das Resultat sollte Sie nicht überraschen, weil Sie es im Prinzip schon kennen. Wenn Sie z. B. die Seitenlänge eines Quadrats verdoppeln, dann vervierfacht sich seine Fläche, wird also mit dem Faktor 2^2 multipliziert. Verdoppeln Sie die Seitenlänge eines Kubus, so verachtfacht sich sein Volumen, d. h. es wird mit 2^3 multipliziert. Das gilt ebenso für Parallelogramme und Spate.

Lösung 27.16. Es ergeben sich

$$\det(-\mathbf{E}_{21}) = \det((-1) \cdot \mathbf{E}_{21}) = (-1)^{21} \cdot \det(\mathbf{E}_{21}) = (-1)^{21} = -1$$

und

$$\det(\mathbf{E}_{10} + \mathbf{E}_{10}) = \det(2 \cdot \mathbf{E}_{10}) = 2^{10} \cdot \det(\mathbf{E}_{10}) = 2^{10} = 1024$$

nach Aufgabe 27.15.

Lösung 27.17. Die Matrix zur Abbildung ist

$$\begin{pmatrix} 1 & 7 \\ 1 & 3 \end{pmatrix}$$

mit der Determinante -4 . Diese Abbildung vergrößert also alle Flächen um den Faktor 4. Da der Einheitskreis die Fläche π hat, hat die Ellipse die Fläche 4π . Siehe dazu die Skizze auf Seite 212.

Lösung 27.18. Das geht im Prinzip wie in Aufgabe 27.4. Die den Spat aufspannenden Vektoren sind $\mathbf{p}_2 - \mathbf{p}_1$, $\mathbf{p}_3 - \mathbf{p}_1$ und $(\mathbf{p}_4 - \mathbf{p}_3) - (\mathbf{p}_2 - \mathbf{p}_1)$. (Der dritte Vektor zeigt vom Punkt P_1 auf die hintere linke Ecke des „Bodens“ des Spats.) Die Matrix, deren Spalten diese drei Vektoren sind, ist

$$\begin{pmatrix} 1 & 1 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 4 \end{pmatrix}$$

mit der Determinante 8, die man sofort ablesen kann. Das ist das gesuchte Volumen.

Lösung 27.20. Die Determinante von $\mathbf{A}$ ist -3 . Nach Lemma 27.2 ist die Determinante von $\mathbf{A}^5$ daher $(-3)^5 = -243$.

Lösung 27.21. Nach Definition der inversen Matrix ist $\mathbf{A} \cdot \mathbf{A}^{-1}$ die Einheitsmatrix $\mathbf{E}_n$ der entsprechenden Größe. Mit Lemma 27.2 und Aufgabe 27.13 erhält man

$$1 = \det(\mathbf{E}_n) = \det(\mathbf{A} \cdot \mathbf{A}^{-1}) = \det(\mathbf{A}) \cdot \det(\mathbf{A}^{-1}).$$

Nun muss man nur noch nach $\det(\mathbf{A}^{-1})$ auflösen und bekommt als Ergebnis

$$\det(\mathbf{A}^{-1}) = 1 / \det(\mathbf{A}) = \det(\mathbf{A})^{-1}.$$

Beachten Sie, dass der Exponent -1 hier zwei unterschiedliche Bedeutungen hat. (Welche?)

Lösung 27.22. Die Determinante von $\mathbf{A}$ ist -12 . Nach Aufgabe 27.21 ist die Determinante von $\mathbf{A}^{-1}$ daher $-1/12$.

Lösung 27.25. Die Determinante von $\mathbf{A}$ ist $3a - 9$, die von $\mathbf{A}^2$ also $(3a - 9)^2$. Damit das 9 wird, muss $3a - 9$ einen der Werte 3 oder -3 annehmen. a muss also 4 oder 2 sein.

Lösung 27.28. Nach Satz 27.3 kann man das anhand der Determinante der zugehörigen Matrix entscheiden. Da diese den Wert 7 hat, ist die Abbildung injektiv.

Lösung 27.29. Die Funktion lässt sich als Komposition aus $(x, y) \mapsto (3x - 2y, 2x)$ und $(x, y) \mapsto (x, y) + (0, 1)$ darstellen. Die erste Abbildung ist linear mit einer nicht verschwindenden Determinante, also injektiv. Die zweite ist eine **Translation**, also sicher auch injektiv. Damit ist auch die Komposition injektiv. (Siehe dazu das Ende von Kapitel IV.)

Lösung 28.1. Man muss das jeweils nur mal hingeschrieben haben, um zu sehen, dass alle drei Eigenschaften wirklich offensichtlich sind. Wir machen das exemplarisch nur für die dritte im zweidimensionalen Fall:

$$\left(\begin{pmatrix} a_1 \\ a_2 \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} \right) \cdot \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} = \begin{pmatrix} a_1 + b_1 \\ a_2 + b_2 \end{pmatrix} \cdot \begin{pmatrix} c_1 \\ c_2 \end{pmatrix}$$

$$\begin{aligned}
&= (a_1 + b_1) \cdot c_1 + (a_2 + b_2) \cdot c_2 \\
&= a_1 c_1 + b_1 c_1 + a_2 c_2 + b_2 c_2 \\
\begin{pmatrix} a_1 \\ a_2 \end{pmatrix} \cdot \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \end{pmatrix} \cdot \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} &= a_1 c_1 + a_2 c_2 + b_1 c_1 + b_2 c_2
\end{aligned}$$

Lösung 28.3. Es ergibt sich:

$$\begin{aligned}
(\mathbf{a} - \mathbf{b})^2 &= (\mathbf{a} - \mathbf{b}) \cdot (\mathbf{a} - \mathbf{b}) \\
&= \mathbf{a} \cdot (\mathbf{a} - \mathbf{b}) - \mathbf{b} \cdot (\mathbf{a} - \mathbf{b}) \\
&= \mathbf{a} \cdot \mathbf{a} - \mathbf{a} \cdot \mathbf{b} - \mathbf{b} \cdot \mathbf{a} + \mathbf{b} \cdot \mathbf{b} \\
&= \mathbf{a}^2 - 2\mathbf{a} \cdot \mathbf{b} + \mathbf{b}^2
\end{aligned}$$

(Machen Sie sich für die einzelnen Zwischenschritte jeweils klar, dass man das wirklich tun darf. Teilweise habe ich zwei Schritte zusammengefasst oder stillschweigend die Kommutativität des Skalarprodukts ausgenutzt.)

Das ist natürlich nichts weiter als die zweite **binomische Formel** für das Skalarprodukt. Es ging bei dieser Aufgabe in erster Linie darum, Vertrautheit mit dem Skalarprodukt zu entwickeln.

Lösung 28.4. Mittels Lemma 28.1 kann man das ganz allgemein ausrechnen:

$$\|\lambda \mathbf{a}\| = \sqrt{(\lambda \cdot \mathbf{a}) \cdot (\lambda \cdot \mathbf{a})} = \sqrt{(\lambda \cdot \lambda) \cdot (\mathbf{a} \cdot \mathbf{a})} = \sqrt{\lambda^2} \cdot \sqrt{\mathbf{a} \cdot \mathbf{a}} = |\lambda| \cdot \|\mathbf{a}\|$$

Beachten Sie dabei, dass $\sqrt{\lambda^2}$ den Wert $|\lambda|$ und nicht etwa λ hat; siehe Aufgabe 12.19.

Lösung 28.5. $\mathbf{a}$ hat die Norm $\sqrt{3^2 + 4^2 + (-1)^2} = \sqrt{26}$. Wir wissen bereits, dass $\lambda \cdot \mathbf{a}$ für jeden Skalar $\lambda \neq 0$ parallel zu $\mathbf{a}$ ist und dass dieser Vektor auch noch in dieselbe Richtung wie $\mathbf{a}$ zeigt, wenn λ positiv ist. Wir müssen also offensichtlich $\mathbf{a}$ einfach mit $1/\sqrt{26}$ multiplizieren, dann erhalten wir mithilfe der vorherigen Aufgabe

$$\left\| \frac{1}{\sqrt{26}} \cdot \mathbf{a} \right\| = \left| \frac{1}{\sqrt{26}} \right| \cdot \|\mathbf{a}\| = \frac{1}{\sqrt{26}} \cdot \sqrt{26} = 1.$$

Ganz allgemein kann man zu jedem Vektor $\mathbf{a}$, der nicht gerade der Nullvektor ist, einen Vektor finden, der in dieselbe Richtung zeigt, aber die Länge 1 hat: man nehme einfach $(1/\|\mathbf{a}\|) \cdot \mathbf{a}$. Man sagt dann auch, dass man $\mathbf{a}$ *normiert* hat.

normieren

Lösung 28.6. Auf Seite 181 kann man erkennen, dass jeder einzelne Eintrag eines Produktes von Matrizen im Prinzip ein Skalarprodukt einer Zeile des ersten Faktors mit einer Spalte des zweiten ist. Im angesprochenen Beispiel ist etwa 8 das Skalarprodukt der zweiten Zeile und der dritten Spalte.

Insbesondere kann man das Skalarprodukt $\mathbf{a} \cdot \mathbf{b}$ auch als das Matrixprodukt $\mathbf{a}^T \cdot \mathbf{b}$ interpretieren. Dabei betrachtet man im zweiten Produkt die Vektoren als Matrizen mit nur einer Spalte. Das Ergebnis ist dann formal eine 1×1 -Matrix, die wir einfach mit ihrem einzigen Eintrag identifizieren.

$$(a_1 \ a_2 \ a_3) \cdot \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix} = (a_1 b_1 + a_2 b_2 + a_3 b_3)$$

Lösung 28.7. Es ergeben sich ungefähr 42° .

Lösung 28.8. Ja, das gilt auch umgekehrt. Das liegt schlicht und einfach daran, dass das Skalarprodukt völlig symmetrisch definiert ist; beide Vektoren werden gleich behandelt.

Lösung 28.9. Bei einem Vektor in der Ebene ist es ganz einfach, einen dazu orthogonalen zu finden. Ist (x, y) orthogonal zu $\mathbf{v}$, so gilt $3x + 7y = 0$ bzw. $3x = -7y$. Das lässt sich natürlich ganz einfach erreichen, indem man $x = 7$ und $y = -3$ wählt oder umgekehrt $x = -7$ und $y = 3$. Man vertauscht also einfach die erste mit der zweiten Komponente und ändert bei einer von beiden das Vorzeichen.³⁰ Skalare Vielfache dieser beiden Vektoren sind natürlich ebenfalls orthogonal zu $\mathbf{v}$, was ja auch geometrisch klar ist.

Falsch wäre es übrigens, als Antwort für diese Aufgabe den Nullvektor zu nennen. Zwar ist das Skalarprodukt von $\mathbf{v}$ und dem Nullvektor tatsächlich null, aber wie wir oben schon bemerkt haben, ergibt es keinen Sinn, über den Winkel zwischen dem Nullvektor und einem anderen Vektor zu sprechen.

Lösung 28.10. Sie sind hoffentlich auf eine Idee wie $(7, -3, 0)$ gekommen. Man kann sich zwei Komponenten auswählen, mit denen wie in der vorherigen Aufgabe verfahren und die dritte Komponente verschwinden lassen. $(2, 0, 3)$ würde es also auch tun. Und übrigens auch jede Linearkombination von so erhaltenen Vektoren. Z. B. ist

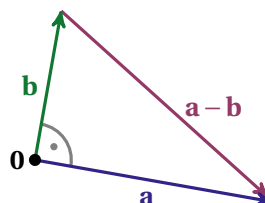
$$3 \cdot \begin{pmatrix} 7 \\ -3 \\ 0 \end{pmatrix} + 5 \cdot \begin{pmatrix} 2 \\ 0 \\ 3 \end{pmatrix} = \begin{pmatrix} 31 \\ -9 \\ 15 \end{pmatrix}$$

auch ein Vektor, der orthogonal zu $\mathbf{v}$ ist. Können Sie mit den Rechengesetzen für das Skalarprodukt begründen, warum das so sein muss? Und ist Ihnen geometrisch klar, warum es im Raum so viele Vektoren gibt, die orthogonal zu einem vorgegeben sind?

Lösung 28.11. Wenn $\mathbf{a}$ und $\mathbf{b}$ orthogonal zueinander sind, ist $\mathbf{a} \cdot \mathbf{b} = 0$. Das Ergebnis von Aufgabe 28.3 vereinfacht sich dann:

$$(\mathbf{a} - \mathbf{b})^2 = \mathbf{a}^2 + \mathbf{b}^2$$

Und das ist natürlich einfach der Satz des Pythagoras. Die Grafik von Seite 223 sieht für orthogonale Vektoren so aus:



³⁰Zur Erinnerung: Das entspricht Drehungen um $\pm 90^\circ$.

Warum lässt sich der Satz mithilfe des Skalarproduktes durch eine einfache algebraische Umstellung beweisen? Weil wir ihn in unsere Definitionen quasi schon „eingebaut“ haben. Dass wir $\mathbf{a}^2$ als das Quadrat der Länge von $\mathbf{a}$ interpretieren, haben wir ja mit dem Satz des Pythagoras begründet.

Wenn die beiden Vektoren nicht senkrecht aufeinander stehen, so entspricht die Darstellung aus Aufgabe 28.3 übrigens dem **Kosinussatz**. Sehen Sie das?

Lösung 28.12. Es ergeben sich $(13, -14, -16)$ und $(-13, 14, 16)$. Das ist kein Zufall. Allgemein gilt $\mathbf{a} \times \mathbf{b} = -(\mathbf{b} \times \mathbf{a})$, wie man leicht nachrechnen kann. Man nennt diese Eigenschaft *Antikommutativität*.

Antikommutativität

Lösung 28.13. Es ergibt sich:

$$\begin{aligned} (\mathbf{a} \times \mathbf{b}) \cdot \mathbf{a} &= (a_2 b_3 - a_3 b_2) a_1 + (a_3 b_1 - a_1 b_3) a_2 + (a_1 b_2 - a_2 b_1) a_3 \\ &= a_1 a_2 b_3 - a_1 a_3 b_2 + a_2 a_3 b_1 - a_1 a_2 b_3 + a_1 a_3 b_2 - a_2 a_3 b_1 = 0 \end{aligned}$$

$\mathbf{a} \times \mathbf{b}$ und $\mathbf{a}$ stehen also senkrecht aufeinander.

Lösung 28.14. Da kommt der Nullvektor heraus. Sind $\mathbf{a}$ und $\mathbf{b}$ nämlich parallel, so kann man $\mathbf{b}$ als $\lambda \mathbf{a}$ darstellen. In Gleichung (28.2) eingesetzt ergibt sich dann:

$$\mathbf{a} \times (\lambda \mathbf{a}) = \begin{pmatrix} a_2 \lambda a_3 - a_3 \lambda a_2 \\ a_3 \lambda a_1 - a_1 \lambda a_3 \\ a_1 \lambda a_2 - a_2 \lambda a_1 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$$

Lösung 28.15. Man berechnet $\mathbf{x} = (4 - 3a, 4, a - 4)$. Bildet man das Skalarprodukt dieses Vektors mit $\mathbf{e}_3$, so ergibt sich $a - 4$. Damit null herauskommt, muss also $a = 4$ gelten.

Lösung 28.17. Man muss nur einsetzen. Der erste Punkt liegt nicht auf h , der zweite schon.

Lösung 28.18. Wie in Aufgabe 28.9 finden wir ganz einfach einen Normalenvektor, z. B. $\mathbf{n} = (5, 2)$. Durch das Skalarprodukt $(5, 2) \cdot (-2, 1)$ erhalten wir -8 , wodurch sich die Normalenform

$$g = \{\mathbf{x} \in \mathbb{R}^2 : (5, 2) \cdot \mathbf{x} + 8 = 0\}$$

ergibt. Wie bei der Punkt-Richtungs-Form gibt es auch hier natürlich unendlich viele Möglichkeiten. Man hätte einen anderen Normalenvektor – z. B. $(-5, -2)$ oder $(10, 4)$ – oder auch einen anderen Punkt statt $(-2, 1)$ nehmen können. Wenn Ihre Lösung also anders als meine aussieht, muss sie deswegen nicht falsch sein.

Lösung 28.19. Steht $\mathbf{n}$ senkrecht auf der Ebene, so sind alle Vektoren, die zwei verschiedene Punkte der Ebene verbinden, orthogonal zu $\mathbf{n}$.

Lösung 28.20. Man erhält (siehe Aufgabe 20.24) als Richtungsvektoren beispielsweise $(2, -3, 4)$ und $(3, -2, -7)$. Deren Vektorprodukt ist $\mathbf{n} = (29, 26, 5)$. Bildet man das Skalarprodukt von $\mathbf{n}$ mit $\mathbf{p}_1$, so erhält man 86. Damit ist

$$E = \{\mathbf{x} \in \mathbb{R}^3 : (29, 26, 5) \cdot \mathbf{x} - 86 = 0\}$$

eine mögliche Normalenform für E .

Lösung 28.22. Durch das Vektorprodukt $\mathbf{n} = (3, -4, 2) \times (-1, -1, 4) = (-14, -14, 7)$ erhält man einen Vektor, der senkrecht auf den beiden vorgegebenen steht. Dessen Länge ist $\|\mathbf{n}\| = 21$. Man muss also mit $1/7$ multiplizieren, um die gewünschte Länge zu erhalten. Damit kann man als Antwort $(-2, -2, 1)$ nehmen, man hätte aber auch $(2, 2, -1)$ wählen können.

Lösung 28.23. Der Term $\mathbf{n} \cdot \mathbf{x} - k$ kann auch negativ sein. (Siehe dazu die Bemerkung zur „Länge der Projektion“ auf Seite 225.) Den tatsächlichen Abstand erhält man also, wenn man den Betrag dieses Ausdrucks nimmt.

Lösung 28.24. Der Normalenvektor $(5, 2)$ hat die Norm $\sqrt{29}$. Also ist $1/\sqrt{29} \cdot (5, 2)$ ein normierter Normalenvektor. Wenn wir den mit dem Ortsvektor $(-2, 1)$ multiplizieren, wird aus dem schon berechneten Wert 8 die Zahl $8/\sqrt{29}$. Damit ergibt sich

$$g = \{\mathbf{x} \in \mathbb{R}^2 : 1/\sqrt{29} \cdot ((5, 2) \cdot \mathbf{x} + 8) = 0\}.$$

Lösung 28.25. Nach unseren Vorüberlegungen müssen wir nur $(4, -3)$ für $\mathbf{x}$ einsetzen. Wir erhalten $22/\sqrt{29}$, also ungefähr 4.1.

Lösung 28.26. Nein, zwei. Ist $\mathbf{n}$ normierter Normalenvektor, dann ist $-\mathbf{n}$ das ebenfalls.

Lösung 28.27. Nein, das Vorzeichen ist falsch. Richtig wäre

$$g = \{\mathbf{x} \in \mathbb{R}^2 : 1/\sqrt{29}((-5, -2) \cdot \mathbf{x} - 8) = 0\}.$$

Lösung 28.28. Das Vektorprodukt der Vektoren $\mathbf{p}_2 - \mathbf{p}_1$ und $\mathbf{p}_3 - \mathbf{p}_1$ ist $\mathbf{n} = (5, -4, 9)$ mit der Norm $\sqrt{122}$. Das Skalarprodukt von $\mathbf{n}$ und $\mathbf{p}_1$ ist 29. Damit erhalten wir

$$E = \{\mathbf{x} \in \mathbb{R}^3 : 1/\sqrt{122} \cdot ((5, -4, 9) \cdot \mathbf{x} - 29) = 0\}.$$

Einsetzen von $\mathbf{q}_1$ bis $\mathbf{q}_3$ in die Gleichung liefert in dieser Reihenfolge die Werte $7/\sqrt{122}$, $-29/\sqrt{122}$ und $3/\sqrt{122}$. Die Antwort ist also $29/\sqrt{122}$.

Lösung 29.1. Wir wissen, dass lineare Abbildungen allgemein Quadrate bzw. Würfel auf Parallelogramme bzw. Spate abbilden. Wenn aber alle Abstände erhalten bleiben, dann müssen aus diesen Quadraten und Würfeln wieder Quadrate und Würfel derselben Größe werden. Also kann die Determinante nur einen der Werte 1 und -1 haben.

Lösung 29.2. $\mathbf{B}$ ist orthogonal, weil $\mathbf{B}^T \cdot \mathbf{B}$ die Einheitsmatrix $\mathbf{E}_3$ ist. $\mathbf{A}$ ist nicht orthogonal, weil $\mathbf{A}^T \cdot \mathbf{A}$ nicht $\mathbf{E}_2$ ist.

Lösung 29.3. Da gibt es viele Möglichkeiten, z. B.

$$\begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}.$$

Entscheidend ist, dass man die Aussage von Aufgabe 29.1 nicht umkehren kann, dass man also wegen $\det(\mathbf{A}) = 1$ nicht schlussfolgern kann, dass $\mathbf{A}$ orthogonal ist.

Lösung 29.4. Das sah bei Ihnen hoffentlich so aus:

$$\begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \cdot \begin{pmatrix} \cos \alpha & \sin \alpha \\ -\sin \alpha & \cos \alpha \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$$

Die Einsen ergeben sich mit dem Satz des Pythagoras, die Nullen, weil sich Terme mit unterschiedlichen Vorzeichen gegenseitig aufheben.

Lösung 29.5. Bei dieser Aufgabe ging es nur darum, ob Sie den vorherigen Teil aufmerksam gelesen haben. Dann ist nämlich ohne große Rechnung klar, dass $\|f(\mathbf{a})\| = \|\mathbf{a}\| = \sqrt{3^2 + 4^2} = 5$ gelten muss.

Lösung 29.6. Multiplikation mit der transponierten Matrix ergibt

$$\frac{1}{\sqrt{2}} \cdot \begin{pmatrix} a & 0 & -1 \\ -1 & 0 & -1 \\ 0 & \sqrt{2} & 0 \end{pmatrix} \cdot \frac{1}{\sqrt{2}} \cdot \begin{pmatrix} a & -1 & 0 \\ 0 & 0 & \sqrt{2} \\ -1 & -1 & 0 \end{pmatrix} = \frac{1}{2} \cdot \begin{pmatrix} a^2 + 1 & 1 - a & 0 \\ 1 - a & 2 & 0 \\ 0 & 0 & 2 \end{pmatrix}.$$

Damit das die Einheitsmatrix $\mathbf{E}_3$ ist, muss $a^2 + 1 = 2$ und $1 - a = 0$ gelten. Das ist nur mit $a = 1$ möglich.

Lösung 29.8. Weil $\mathbf{A}$ eine Drehung beschreibt, gilt $\det \mathbf{A} = 1$. Es folgt

$$\begin{aligned} \det(\mathbf{A}^{-1} \cdot \mathbf{B}^{-1} \cdot \mathbf{C}^T) &= \det(\mathbf{A}^{-1}) \cdot \det(\mathbf{B}^{-1}) \cdot \det(\mathbf{C}^T) \\ &= \det(\mathbf{A})^{-1} \cdot \det(\mathbf{B})^{-1} \cdot \det(\mathbf{C}) = 1 \cdot \frac{1}{6} \cdot 24 = 4. \end{aligned}$$

Lösung 29.9. Man erhält $vu = -27 - 8i + 19j + 17k$ und das ist leider weder uv noch $-uv$. Im Allgemeinen darf man also *nicht* uv durch vu ersetzen, wenn u und v beliebige Quaternionen sind. Allerdings gilt offensichtlich $uv = vu$, wenn eine der beiden Zahlen reell ist.

Lösung 29.11. Es ergeben sich $\bar{u} = 3 + 2i - 5j + k$, $|u| = \sqrt{39}$ und

$$u^{-1} = \frac{1}{39} \cdot (3 + 2i - 5j + k) = \frac{1}{13} + \frac{2}{39} \cdot i - \frac{5}{39} \cdot j + \frac{1}{39} \cdot k.$$

Lösung 29.13. Es ergeben sich mit $1/39 \cdot (12 - 5i - j + 8k)$ und $1/13 \cdot (4 + 3i + j)$ auch hier zwei ganz unterschiedliche Werte. Trotz der Existenz von Kehrwerten ist es daher nicht sinnvoll, eine Division wie „ v/u “ einzuführen, weil nicht klar wäre, welches von beiden Produkten gemeint ist. (Ein Term wie v/a ist allerdings OK, wenn a eine reelle Zahl ist, weil $v \cdot a^{-1} = a^{-1} \cdot v$ gilt.)

Lösung 29.14. Es ergibt sich $t = 1/(5\sqrt{2}) \cdot (5 + 3i - 4k)$. Der gedrehte Vektor ist $1/25 \cdot (-62, -115, 41)$.

Lösung 30.1. Das sind 2, 4, 16, 64 und 1024.

Lösung 30.2. Zum Beispiel $a_n = 2n/(2n+1)$. Es gibt in solchen Fällen jedoch nie nur eine Lösung. Man könnte auch

$$a_n = -\frac{16n^4}{945} + \frac{8n^3}{45} - \frac{92n^2}{135} + \frac{374n}{315}$$

nehmen. Das wäre zwar ziemlich haarsträubend, aber es würde die Aufgabe auch lösen. (Rechnen Sie nach!) Siehe dazu auch das nicht ganz ernst gemeinte Video *Wie man JEDEN Intelligenztest besteht*.

Lösung 30.3. Ja. Die Menge der „Ausnahmen“ ist leer und die leere Menge ist natürlich endlich.

Lösung 30.4. Nein. Zwar sind unendlich viele Folgenglieder gerade, aber die Menge der „Ausnahmen“ ist die unendliche Menge der ungeraden Zahlen. *Fast alle* ist ein stärkerer Begriff als *unendlich*.

Lösung 30.6. Der „Trick“ ist eigentlich immer, den jeweiligen Term so umzuformen, dass man Summe, Differenz, Produkt oder Quotient von Folgen erhält, von denen man bereits weiß, gegen was sie konvergieren. In diesem Fall bieten sich

$$\frac{1}{n^3} = \frac{1}{n^2} \cdot \frac{1}{n}, \quad \frac{5}{n} = 5 \cdot \frac{1}{n} \quad \text{und} \quad \frac{2n-1}{n} = 2 - \frac{1}{n}$$

an und man erhält die Grenzwerte $0 \cdot 0 = 0$, $5 \cdot 0 = 0$ und $2 - 0 = 2$.

Lösung 30.7. Nein, die Folge ist divergent. Offensichtlich liegen in den beiden disjunkten Intervallen $(-3/2, -1/2)$ und $(1/2, 3/2)$ jeweils unendlich viele Folgenglieder. Daher kann weder 1 noch -1 der Grenzwert sein und eine andere Zahl kommt „erst recht“ nicht infrage.

Lösung 30.8. Wenn der Grenzwert b der Folge (b_n) nicht null ist, dann haben fast alle Folgenglieder dieser Folge einen Abstand von b , der kleiner als $|b|$ ist, und können daher nicht die Zahl 0 sein. Die Null kann also nur endlich oft als Folgenglied in (b_n) vorkommen.

Lösung 30.9. Die erste Folge geht gegen $-1/2$ (man beachte das Vorzeichen). Hier wird wie gesagt mit $1/n^2$ erweitert. Und durch Erweitern mit $1/n^3$ ergibt sich bei der zweiten Folge der Grenzwert 0. Sie erkennen für Folgen dieser Art hoffentlich ein allgemeines Verfahren, bei dem man immer mit dem Kehrwert der höchsten Potenz von n erweitert, wenn in Zähler und Nenner **Polynome** in n stehen. Mit etwas Übung sieht man den Grenzwert, ohne dass man die Berechnung explizit durchführen muss.

Lösung 30.10. Mit dem besagten „Trick“ (Erweitern mit n^{-3}) erhält man den Grenzwert $6/a$. Damit daraus 42 wird, muss $a = 1/7$ gelten.

Lösung 30.11. Im Zähler ist der dominierende Term $12n^7$. Damit der Nenner „gewinnt“, muss n dort mindestens in der achten Potenz stehen, also muss k mindestens 4 sein.

Lösung 30.12. Wie wäre es mit $(1 + 1/n)$? Es gibt natürlich viele Beispiele.

Lösung 30.13. Hier muss man ein bisschen umformen und dann wieder mit dem Kehrwert des am schnellsten wachsenden Terms erweitern. Es ergibt sich:

$$\frac{4^{n+1} + 2^n}{4^n + 2^{2n}} = \frac{4^{n+1} + 2^n}{2 \cdot 4^n} \cdot \frac{4^{-n}}{4^{-n}} = \frac{4 + (1/2)^n}{2} \xrightarrow{n \rightarrow \infty} \frac{4 + 0}{2} = 2$$

Lösung 30.14. Man muss lediglich eine einfache Umformung erkennen, um auf eine Folge zu kommen, deren Grenzwert wir gerade hergeleitet haben:

$$\sqrt[n]{2^{n+1}} = \sqrt[n]{2^n \cdot 2} = \sqrt[n]{2^n} \cdot \sqrt[n]{2} = 2 \cdot \sqrt[n]{2} \rightarrow 2 \cdot 1 = 2$$

Lösung 30.15. Sie sollten beobachtet haben, dass die Folgenglieder sehr schnell sehr groß werden. Steht das nicht im Widerspruch zur Aussage von Satz 30.2? Nein. Die Folge konvergiert tatsächlich gegen null, aber man kann diesen Effekt erst beobachten, wenn man für n entsprechend große Werte einsetzt. Die „Moral“ dieser Geschichte ist, dass man, wenn es um Grenzwerte geht, nicht einfach die ersten zehn, hundert oder tausend Folgenglieder berechnen und daraus Schlüsse ziehen kann.

Lösung 30.17. Die alternierende Folge aus Aufgabe 30.7 könnte man nehmen.

Lösung 30.18. Zum Beispiel $((-1)^n/n)$.

Offensichtlich gilt die Umkehrung von Lemma 30.4 jedoch, wenn (a_n) eine reelle Nullfolge ist, die nie das Vorzeichen wechselt: Dann ist $(1/a_n)$ bestimmt divergent.

Lösung 31.1. Für ein einzelnes Paar aus einer Zeile der ersten und einer Spalte der zweiten Matrix sind es n Multiplikationen und $n - 1$ Additionen, also $2n - 1$ Operationen. Da man das für jede der n^2 möglichen Kombinationen aus Zeilen und Spalten machen muss, ergeben sich insgesamt $n^2(2n - 1) = 2n^3 - n^2$ Operationen.

Lösung 31.2. In den meisten Fällen kann man Lemma 31.2 anwenden. Die Folge $(2^{n+42}/3^n)$ ist die Nullfolge $(2^{42} \cdot (2/3)^n)$ und daher ist die erste Aussage wahr und die zweite nicht. $(2^n/(n^2 2^n))$ ist die Nullfolge $(1/n^2)$ und daher ist die dritte Aussage falsch. $(n^2 2^n/3^n)$ ist die Folge $(n^2/(3/2)^n)$, die nach Satz 30.2 eine Nullfolge ist, also ist die vierte Aussage wahr. $(2^n/(2^n)^2)$ ist die Nullfolge $(1/2^n)$ und daher ist die fünfte Aussage falsch. Dass die letzte Aussage richtig ist, folgt schließlich mit $\log n \in \mathcal{O}(n)$ und Punkt (iii) von Satz 31.1.

Lösung 32.1. a_k ist in diesem Fall k^2 . Es müsste daher so aussehen:

$$\sum_{k=3}^5 k^2 = 3^2 + 4^2 + 5^2 = 50$$

Lösung 32.2. Das sollte so aussehen:

$$\begin{aligned} s_1 &= \sum_{k=1}^{101} 4k & s_2 &= \sum_{i=3}^{1000} (3i + 1) \\ s_3 &= \sum_{n=1}^{11} n^{-2} = \sum_{n=1}^{11} \frac{1}{n^2} & s_4 &= \sum_{j=1}^{42} (-1)^{j+1} j \\ s_5 &= \sum_{k=2}^{33} n a_{3k} \end{aligned}$$

Beachten Sie, dass bei solchen Summen die „Schrittweite“ immer 1 ist. Um z. B. in s_1 Schritte der Weite 4 zu erreichen, werden die Summanden entsprechend multipliziert. Bei s_2 wird außerdem der konstante Term 1 addiert, um das „Verschieben“

der Summe zu erreichen. Bei s_4 sorgt die Potenz von -1 für den Vorzeichenwechsel. Es wurden absichtlich unterschiedliche Namen für die Laufvariablen verwendet, um klarzumachen, dass diese Namen keine Rolle spielen. (Sofern sie nicht in den Summanden auftauchen. Bei s_5 sollte die Laufvariable also *nicht* n heißen.)

Lösung 32.4. $1 \cdot 1 + 1 \cdot 1$ ist nicht dasselbe wie $(1 + 1) \cdot (1 + 1)$. Man kann als Gegenbeispiel also $a_k = b_k = 1$, $m = 1$ und $n = 2$ nehmen.

Lösung 32.5. Bei der ersten Summe kann man zunächst den Faktor 3 ausklammern und erkennt dann eine arithmetische Summe, die „zu spät anfängt“ und daher in zwei Teile zerlegt werden muss. Bei der zweiten Summe handelt es sich im Wesentlichen um eine geometrische Summe mit $q = 1/2$.

$$\begin{aligned}\sum_{n=5}^{20} 3n &= 3 \sum_{n=5}^{20} n = 3 \left(\sum_{n=1}^{20} n - \sum_{n=1}^4 n \right) = 3 \left(\frac{20 \cdot 21}{2} - \frac{4 \cdot 5}{2} \right) = 600 \\ \sum_{n=3}^8 \frac{1}{2^n} &= \sum_{n=0}^8 \left(\frac{1}{2} \right)^n - \sum_{n=0}^2 \left(\frac{1}{2} \right)^n = \frac{\left(\frac{1}{2} \right)^9 - 1}{\frac{1}{2} - 1} - \frac{\left(\frac{1}{2} \right)^3 - 1}{\frac{1}{2} - 1} = 2 \left(\left(\frac{1}{8} - 1 \right) - \left(\frac{1}{512} - 1 \right) \right) \\ &= \frac{1}{4} - \frac{1}{256} = \frac{63}{256}\end{aligned}$$

Lösung 32.6. Die „Schrittweite“ ist hier offenbar 10 statt 1, also klammert man 10 aus und bekommt

$$\begin{aligned}40 + 50 + 60 + \dots + 230 &= 10(4 + 5 + 6 + \dots + 23) = 10 \sum_{k=4}^{23} k = 10 \left(\sum_{k=1}^{23} k - \sum_{k=1}^3 k \right) \\ &= 10 \left(\frac{23 \cdot 24}{2} - \frac{3 \cdot 4}{2} \right) = 2700.\end{aligned}$$

Lösung 32.7. Hier kann man im Prinzip wie in Aufgabe 32.6 vorgehen (die Schrittweite ist offenbar 3), muss aber beachten, dass die Summanden keine Vielfachen von 3 sind und deshalb „verschoben“ werden müssen:

$$\begin{aligned}7 + 10 + 13 + 16 + \dots + 1000 &= (1 + 6) + (1 + 9) + (1 + 12) + \dots + (1 + 999) \\ &= \sum_{k=2}^{333} (1 + 3k) = \left(\sum_{k=2}^{333} 1 \right) + 3 \sum_{k=2}^{333} k \\ &= 332 \cdot 1 + 3 \cdot \left(\frac{333 \cdot 334}{2} - 1 \right) = 167\,162\end{aligned}$$

Lösung 32.8. Mit $\sum_{k=1}^n k^2 = n^3/3 + n^2/2 + n/6$ ergibt sich

$$\sum_{n=1}^8 n^2 = \frac{1}{3} \cdot 8^3 + \frac{1}{2} \cdot 8^2 + \frac{1}{6} \cdot 8 = \frac{512 + 96 + 4}{3} = 204.$$

Für die zweite Summe erhält man 9450. Hier ist natürlich zu beachten, dass die Summe erst bei 3^2 und nicht bei 1^2 beginnt.

Lösung 32.9. Als Reihe hat das Beispiel die Form $\sum_{k=0}^{\infty} (-1)^k$ und die Folge der Partialsummen ist $1, 0, 1, 0, 1, 0, \dots$. Da diese Folge nicht konvergiert, gibt es keinen Reihenwert.

Lösung 32.10. Mit $q = 1/3$ erhalten wir $\sum_{n=0}^{\infty} 3^{-n} = (1 - 1/3)^{-1} = 3/2$. Und mit $\sum_{n=0}^1 3^{-n} = 1 + 1/3 = 4/3$ ergibt sich als gesuchter Reihenwert $3/2 - 4/3 = 1/6$.

Lösung 32.11. Zum Beispiel so:

$$\begin{aligned}\sum_{n=3}^{\infty} \frac{3}{n+n^2} &= 3 \left(\sum_{n=1}^{\infty} \frac{1}{n+n^2} - \sum_{n=1}^2 \frac{1}{n+n^2} \right) = 3 \left(1 - \frac{1}{2} - \frac{1}{6} \right) = 1 \\ \sum_{n=2}^{\infty} \frac{5}{2^n} &= 5 \left(\sum_{n=0}^{\infty} \frac{1}{2^n} - \sum_{n=0}^1 \frac{1}{2^n} \right) = 5 \left(2 - 1 - \frac{1}{2} \right) = \frac{5}{2} \\ \sum_{n=0}^{\infty} \frac{2 \cdot 4^n + 3^{n+1}}{6^n} &= 2 \sum_{n=0}^{\infty} \left(\frac{2}{3} \right)^n + 3 \sum_{n=0}^{\infty} \left(\frac{1}{2} \right)^n = 2 \left(1 - \frac{2}{3} \right)^{-1} + 3 \left(1 - \frac{1}{2} \right)^{-1} = 12\end{aligned}$$

Lösung 33.1. Das ist offenbar $1/(n+1)$.

Lösung 33.2. Offensichtlich (siehe Aufgabe 33.1) ist die Folge $(n!/(n+1)!)$ (bis auf eine „Verschiebung“) die **harmonische Folge**. Mit Lemma 31.2 und Lemma 33.1 folgt daher sofort, dass die erste Aussage wahr ist und die zweite nicht.

Lösung 33.3. $\exp(0)$ ist 1, weil in diesem Fall alle Reihenglieder bis auf das erste verschwinden. Für e erhält man

$$e = \exp(1) \approx 1 + 1 + \frac{1}{2} + \frac{1}{6} + \frac{1}{24} = 65/24 \approx 2.70833.$$

Wenn man wie hier nach k Summanden abbricht, dann gilt für den fehlenden Teil der Reihensumme

$$\begin{aligned}\sum_{n=k}^{\infty} \frac{1}{n!} &= \frac{1}{k!} \sum_{n=k}^{\infty} \frac{k!}{n!} = \frac{1}{k!} \sum_{n=k}^{\infty} \frac{1}{(k+1)(k+2) \cdots n} \\ &\leq \frac{1}{k!} \sum_{n=0}^{\infty} \frac{1}{(k+1)^n} = \frac{1}{k!} \cdot \frac{1}{1 - 1/(k+1)} = \frac{1}{k!} \cdot \frac{k+1}{k}.\end{aligned}$$

Im konkreten Fall $k = 5$ ist der Fehler also höchstens $1/100$ und daher wäre 2.7 eine sichere **Rundung**. (Der tatsächliche Wert von e auf zehn Stellen nach dem Komma gerundet ist 2.718281828. Da diese Zahl irrational ist, kann man immer nur Näherungswerte angeben.)

Lösung 33.4. Wenn z nicht reell ist, könnte $\exp(z)$ **echt komplex** sein. Was ein beliebiger rationaler Exponent bedeuten soll, haben wir aber nicht definiert. Siehe die Diskussion auf Seite 115.

Lösung 33.5. Einer der beiden Faktoren muss verschwinden, damit sich null ergibt. Da nach Satz 33.2 der zweite Faktor aber niemals verschwindet, kann nur null herauskommen, wenn $2x - 6 = 0$ gilt. Die einzig mögliche Lösung ist also $x = 3$.

Lösung 33.6. Ja. Wer einfach nur 1000 Euro einzahlt und am Ende des Jahres auf sein Konto schaut, wird dort 2000 Euro vorfinden. Sie jedoch haben nach einem halben Jahr 1500 Euro und nach einem weiteren halben Jahr $1.5 \cdot 1500 = 2250$ Euro.

Lösung 33.7. Der Betrag am Ende ist $A \cdot (1 + 1/n)^n$, wobei A die ursprüngliche Einlage (z. B. $A = 1000$) ist. Wenn n gegen unendlich geht, geht der Jahresendbetrag

nach Satz 33.3 gegen $e \cdot A$. Man bekommt zwar umso mehr Geld, je öfter man abhebt und wieder einzahlt, aber es gibt eine Obergrenze. Bei einem Startkapital von 1000 Euro würde man selbst dann, wenn die Bank einem gestatten würde, jede Sekunde abzuheben und wieder einzuzahlen, nicht über 2718 Euro und 28 Cent hinauskommen.

Lösung 33.8. Die Umformung $(1 + 1/n)^{n+1} = (1 + 1/n)^n \cdot (1 + 1/n)$ zeigt, dass die erste Folge nach Lemma 30.1 und Satz 33.3 gegen $e \cdot (1 + 0) = e$ konvergiert. Und das würde ebenso gelten, wenn man den Exponenten $n + 1$ durch $n + 2$ oder $n + 42$ ersetzen würde. Die Umformung $(1 + 1/n)^{2n} = (1 + 1/n)^n \cdot (1 + 1/n)^n$ zeigt, dass die zweite Folge gegen $e \cdot e = e^2$ konvergiert.

Lösung 33.9. Der Kehrwert von $n/(n + 1)$ ist $(n + 1)/n = 1 + 1/n$. Nach Satz 33.3 konvergiert die Folge der Kehrwerte somit gegen e . Also konvergiert die Folge, um die es eigentlich geht, nach Lemma 30.1 gegen $1/e$.

Lösung 33.10. Man sollte im Zähler die **binomische Formel** erkennen. Dann erhält man

$$\left(\frac{n^2}{n^2 + 2n + 1} \right)^n = \left(\frac{n^2}{(n + 1)^2} \right)^n = \left(\frac{n}{n + 1} \right)^n \cdot \left(\frac{n}{n + 1} \right)^n.$$

Mit Aufgabe 33.9 ergibt sich als Grenzwert e^{-2} .

Lösung 33.11. Ihnen ist hoffentlich klar, dass $2^6 = 64$ und $10^5 = 100\,000$ gilt. Daher sind die Antworten 6 und 5.

Lösung 33.12. Man wendet Lemma 33.6 an und braucht sogar nur eine der beiden Funktionen. Entweder $\log_{10} 42 / \log_{10} 7$ oder $\ln 42 / \ln 7$.

Lösung 33.13. Es ergibt sich $a = 1/\sqrt{500}$, $b = \sqrt{5} \cdot \sqrt{500} = \sqrt{2500} = 50$ und schließlich $c = \log_{10} 20b = \log_{10} 1000 = 3$.

Lösung 33.14. Nach Definition gilt $x = c^b$. Nach den bekannten Rechenregeln aus Lemma 33.5 folgt $x^{-2/b} = (c^b)^{-2/b} = c^{-2} = 1/c^2$. Welchen der beiden letzten Ausdrücke man als einfacher oder schöner ansieht, ist Geschmackssache.

Lösung 33.15. Wendet man auf beide Seiten $\log_a$ an, so erhält man mit Lemma 33.6

$$x = \log_a(5/a^2) = \log_a 5 - \log_a a^2 = (\log_a 5) - 2.$$

Alternativ kann man auch erst die Gleichung mit a^2 multiplizieren und dann den Logarithmus anwenden. Das führt über $a^{x+2} = 5$ zu $x + 2 = \log_a 5$, also natürlich zum selben Ergebnis.

Lösung 33.16. Mit $a = 2$ und $b = 8$ erhält man $\log_2 2^8 = 8$, aber $2\log_2 8 = 6$. Die korrekte Formel steht am Ende von Lemma 33.6.

Lösung 34.2. Dass das Produkt von Polynomen wieder ein Polynom ist, sieht man, wenn man ausmultipliziert, z. B.

$$(3x^2 + 2x - 5) \cdot (3x + 1) = 9x^3 + 9x^2 - 13x - 5.$$

Damit sollte auch ersichtlich sein, dass der Grad des Produktes die Summe der Grade der Faktoren ist.

Lösung 34.3. Zum Beispiel $p(x) = 2x^2 + x$ und $q(x) = -2x^2 + x$.

Lösung 34.4. Laut Aufgabenstellung kennen wir diese Werte:

x	1/5	1/4	1/3	1/2	1
$f(x)$	-32	16	-8	4	-2

Das ist allerdings schon mehr Information als nötig. Eigentlich brauchen wir nur das hier:

x	1/5	1/4	1/3	1/2	1
$f(x)$	-	+	-	+	-

Dabei steht – für $f(x) < 0$ und + für $f(x) > 0$. Da f nach Voraussetzung stetig ist, muss es nach dem **Zwischenwertsatz** zwischen + und – bzw. zwischen – und + jeweils mindestens eine Nullstelle geben. Da es vier solche Vorzeichenwechsel gibt, muss f mindestens vier Nullstellen haben.

Lösung 34.5. Es gilt $h(3) = f(3) \cdot g(3) > 0$ und nach Satz 34.1 ist h stetig. Aus $g(4) < 0$ folgt $h(4) < 0$ und nach dem **Zwischenwertsatz** hat h zwischen 3 und 4 eine Nullstelle. Aus keiner der anderen Voraussetzungen kann man mit Sicherheit $h(4) < 0$ schließen.

Lösung 34.7. Der Abstand von a_0 und b_0 ist 4 und durch das fortgesetzte Halbieren ist der Abstand von a_k und b_k somit $2^{-k} \cdot 4 = 2^{2-k}$. Nach dem 15. Schritt ist der Abstand 2^{-13} . (Man hätte natürlich auch einfach $b_{15} - a_{15}$ berechnen können, aber so hat man eine allgemeine Formel.) Die Mitte des Intervalls ist vom richtigen Ergebnis also maximal 2^{-14} entfernt. Das ist ungefähr $6.1 \cdot 10^{-5}$ und damit weniger als $5 \cdot 10^{-4}$, aber mehr als $5 \cdot 10^{-5}$. Man darf also maximal drei Nachkommastellen angeben: -1.167.

Lösung 34.8. Die korrekte Antwort ist 1.40 – und nicht etwa 1.4!

Lösung 34.9. Eine Lösung in PYTHON könnte so aussehen:

```
def zero(f,a,b,n):
    eps = 10**(-n)
    m = (a+b)/2
    while abs(a-m) > eps:
        if f(m) == 0:
            return round(m,n)
        if f(a)*f(m) < 0:
            b = m
        else:
            a = m
        m = (a+b)/2
    return round(m,n)
```

```
zero(lambda x: x**5-x+1,-2,2,3)
```

Lösung 35.1. Sei $f(x) = a$ und $g(x) = x$. Dann ergibt sich

$$f'(x_0) = \lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0} = \lim_{x \rightarrow x_0} \frac{a - a}{x - x_0} = 0 \quad \text{und}$$

$$g'(x_0) = \lim_{x \rightarrow x_0} \frac{g(x) - g(x_0)}{x - x_0} = \lim_{x \rightarrow x_0} \frac{x - x_0}{x - x_0} = 1.$$

Lösung 35.2. Man erhält $2x_0$, und zwar so:

$$\lim_{x \rightarrow x_0} \frac{x^2 - x_0^2}{x - x_0} = \lim_{x \rightarrow x_0} \frac{(x - x_0)(x + x_0)}{x - x_0} = \lim_{x \rightarrow x_0} (x + x_0) = 2x_0$$

Diese Rechnung wurde am Anfang von Abschnitt 34 schon einmal durchgeführt.

Lösung 35.3. Da 1 die Ableitung von $x \mapsto x$ ist, ergibt sich $1 \cdot x_0 + x_0 \cdot 1 = 2x_0$. Natürlich muss dasselbe herauskommen, was wir schon einmal berechnet hatten.

Lösung 35.5. Die Ableitung ist $9x_0^2 - 10x_0 + 8$.

Lösung 35.6. So geht es los:

$$\begin{aligned} f(x) &= 3x^3 - 5x^2 + 8x - 2 \\ f^{(1)}(x) &= f'(x) = 9x^2 - 10x + 8 \\ f^{(2)}(x) &= f''(x) = 18x - 10 \\ f^{(3)}(x) &= 18 \\ f^{(4)}(x) &= 0 \end{aligned}$$

Alle weiteren Ableitungen sind ebenfalls die konstante Nullfunktion.

Lösung 35.7. Mit der Leibnizregel erhält man

$$\begin{aligned} f'(x) &= (2x)(x^3 - x) + (x^2 + 3)(3x^2 - 1) = 2x^4 - 2x^2 + 3x^4 + 9x^2 - x^2 - 3 \\ &= 5x^4 + 6x^2 - 3. \end{aligned}$$

Die alternative Berechnung liefert

$$\begin{aligned} f(x) &= (x^2 + 3)(x^3 - x) = x^5 + 3x^3 - x^3 - 3x = x^5 + 2x^3 - 3x \quad \text{und} \\ f'(x) &= 5x^4 + 6x^2 - 3. \end{aligned}$$

Natürlich muss in beiden Fällen dasselbe herauskommen!

Lösung 35.8. x^{-3} ist der Kehrwert von x^3 und die Ableitung von x^3 ist (nach Aufgabe 35.4) $3x^2$. Damit ergibt sich

$$f'(x) = \frac{-3x^2}{(x^3)^2} = -3x^{-4}.$$

Lösung 35.9. Man muss das nur noch für negative Zahlen beweisen. Sei also $n \in \mathbb{N}^+$. Wir wissen schon, dass nx^{n-1} die Ableitung von x^n ist. Für $f(x) = x^{-n}$ gilt dann

$$f'(x) = \frac{-nx^{n-1}}{(x^n)^2} = \frac{-nx^{n-1}}{x^{2n}} = -nx^{n-1-2n} = -nx^{-n-1}.$$

Lösung 35.10. Die Ableitung des Zählers ist $2x + 2$, die des Nenner $6x$. Damit ergibt sich insgesamt

$$f'(x) = \frac{(2x+2)(3x^2-2) - (x^2+2x+2)6x}{(3x^2-2)^2} = -\frac{6x^2+16x+4}{9x^4-12x^2+4}.$$

Lösung 35.12. Wir haben hier $h = f \circ g$ mit $g(x) = 3x^2 - 5x + 2$ und $f(x) = x^4$. Die „äußere“ Ableitung ist $f'(x) = 4x^3$, die „innere“ $g'(x) = 6x - 5$. Damit ergibt sich

$$h'(x) = 4(3x^2 - 5x + 2)^3(6x - 5) = (3x^2 - 5x + 2)^3(24x - 20).$$

Lösung 35.13. In beiden Fällen muss man die Kettenregel verwenden. Im ersten Fall erhält man $2x \exp(x^2)$. Im zweiten Fall werden die Rollen von „äußerer“ und „innerer“ Funktion vertauscht und man erhält $2 \exp(x)^2$.

Lösung 35.14. Beispiele wären $2 \exp(x)$ oder $\exp(x + 42)$. Generell ist jede Funktion der Form $x \mapsto a \exp(x + b)$ mit Konstanten a und b ihre eigene Ableitung.

Lösung 35.15. Wieder ein Fall für die Kettenregel. Es ergeben sich $2x/(x^2 + 1)$ und $(2 \ln x)/x$.

Lösung 35.16. Das sind $2x \cos(x^2 + 1)$ und $\cos x / \sin x$. Die zweite Funktion wird auch als **Kotangens** bezeichnet.

Lösung 35.17. Die Ableitungsfunktion kann man mit der **Leibnizregel** oder mit der Kettenregel berechnen. In beiden Fällen ergibt sich $f'(x) = 2 \sin x \cos x$. Mit einem Taschenrechner ergibt sich gerundet $f'(6/5) \approx 0.68$. Man hat nun ein Steigungsdreieck mit der waagerechten Kathete 1 und der senkrechten Kathete $f'(6/5)$. Mit dem **Arkustangens** kommt man so auf den Steigungswinkel

$$\varphi = \arctan(f'(6/5)) \approx 0.59.$$

Hier kann man übrigens beobachten, dass man nicht zu früh runden sollte. Setzt man im letzten Schritt für $f'(6/5)$ den gerundeten Wert 0.68 ein, so erhält man – gerundet auf zwei Stellen nach dem Komma – $\arctan(0.68) \approx 0.60$.

Lösung 35.18. Nach der Kettenregel ist die Ableitungsfunktion die Funktion, die x auf $(2x+1)(x^2+x-12)^5$ abbildet. Da ein Produkt nur verschwindet, wenn einer der Faktoren verschwindet, ist eine Nullstelle dieser Funktion entweder eine Lösung von $2x+1=0$ oder eine von $x^2+x-12=0$. Die Antwort ist daher $\{3, -4, -1/2\}$.

Lösung 36.1. Im ersten Fall ergibt sich ein Rechteck mit den gegenüberliegenden Ecken $(a, 0)$ und $(b, 3)$ und damit der Flächeninhalt $3(b-a)$. Im zweiten Fall bekommt man ein Dreieck mit den Ecken $(a, 0)$, $(b, 0)$ und $(b, b-a)$ und damit den Flächeninhalt $(b-a)^2/2$.

Lösung 36.3. Da muss natürlich -2 herauskommen. Entscheidend ist hier das Vorzeichen. Wenn man es genau formuliert, dann ist das Integral weder die Fläche *unter* dem Funktionsgraphen noch die Fläche *zwischen* Funktionsgraph und horizontaler Achse, sondern die *gerichtete Flächenbilanz* zwischen Funktionsgraph und horizontaler Achse: Teile oberhalb der Achse werden dabei positiv gerechnet, Teile unterhalb der Achse negativ.

Lösung 36.4. Weil die Teile oberhalb und unterhalb der horizontalen Achse offenbar gleich groß sind, muss sich als „Bilanz“ null ergeben.

Lösung 36.5. Das funktioniert wegen der Festsetzungen *am Ende der Definition des Integrals*.

Lösung 36.6. Wir wissen, dass $-\sin x$ die Ableitung von $\cos x$ ist, also ist $\sin x$ die Ableitung von $-\cos x$ bzw. $-\cos x$ eine Stammfunktion von $\sin x$. Eine weitere Stammfunktion wäre z.B. $\pi - \cos x$.

Lösung 36.7. Da $(s+1)x^s$ die Ableitung von x^{s+1} ist, müssen wir einfach durch $s+1$ dividieren: $x^{s+1}/(s+1)$ ist eine Stammfunktion von x^s .

Das kann aber offenbar nicht für $s = -1$ gelten, weil man dann durch null teilen müsste. Wir wissen aber schon, dass $1/x$ die Ableitung von $\ln x$ ist. Die Stammfunktion von x^{-1} ist also der natürliche Logarithmus.

Lösung 36.8. Für jedes x ergibt sich die gerichtete Flächenbilanz von a bis x . Man kann sich die rechte Integralgrenze x wie einen Schieberegler vorstellen, während die linke Grenze a fest bleibt. Insbesondere gilt natürlich $F(a) = 0$.

Lösung 36.10. Zunächst brauchen wir Stammfunktionen. Mit Aufgabe 36.7 erhalten wir $x^3/3 + x^2/2 + x$ als eine Stammfunktion für den ersten Integranden. Weil $2^x \ln 2$ die Ableitung von 2^x ist, ist $2^x/\ln 2$ eine Stammfunktion von 2^x . Damit erhalten wir

$$\int_1^3 (x^2 + x + 1) dx = x^3/3 + x^2/2 + x \Big|_1^3 = 44/3 \quad \text{und} \\ \int_0^2 2^x dx = 2^x/\ln 2 \Big|_0^2 = 3/\ln 2.$$

Lösung 36.11. Die Stammfunktionen sind $\ln x$ und $-\cos x$. Damit ergibt sich

$$\int_1^3 \frac{dx}{x} = \ln x \Big|_1^3 = \ln 3 \quad \text{und} \\ \int_{-\pi}^{\pi} \sin x dx = -\cos x \Big|_{-\pi}^{\pi} = 0.$$

Dass das zweite Integral verschwindet, hätte man sich ohne Rechnen überlegen können. Das liegt daran, dass die Integralgrenzen symmetrisch zur Null liegen und der Sinus eine *ungerade* Funktion ist. (Machen Sie sich eine Skizze, wenn Ihnen das nicht klar ist!)

Lösung 36.12. Die Ableitung von $\ln \sin x$ ist nach der Kettenregel $1/\tan x$. Also erhalten wir

$$\int_{\pi/4}^{\pi/2} \frac{2dx}{\tan x} = 2 \int_{\pi/4}^{\pi/2} \frac{dx}{\tan x} = 2 \ln \sin x \Big|_{\pi/4}^{\pi/2} = 2 \ln 1 - 2 \ln(1/\sqrt{2}) = \ln 2.$$

Lösung 37.1. Das könnte z. B. so aussehen:

$$A = \{(a, b) \in \Omega : a \bmod 2 = 0\}$$

$$B = \{(a, b) \in \Omega : b > 3\}.$$

Beide Mengen haben die Mächtigkeit 18.

Lösung 37.2. Die Mengen sind die folgenden:

$$A \cap B = \{(a, b) \in \Omega : a \bmod 2 = 0 \wedge b > 3\}$$

$$= \{(2, 4), (2, 5), (2, 6), (4, 4), (4, 5), (4, 6), (6, 4), (6, 5), (6, 6)\}$$

$$A \cup B = \{(a, b) \in \Omega : a \bmod 2 = 0 \vee b > 3\}$$

$$= \Omega \setminus \{(1, 1), (1, 2), (1, 3), (3, 1), (3, 2), (3, 3), (5, 1), (5, 2), (5, 3)\}$$

$$A^c = \{(a, b) \in \Omega : a \bmod 2 \neq 0\} = \{(a, b) \in \Omega : a \bmod 2 = 1\}.$$

$A \cap B$ ist das Ereignis, dass man sowohl im ersten Wurf eine gerade Augenzahl als auch im zweiten eine größere Augenzahl als 3 wirft. Die Menge hat neun Elemente. Es gibt also neun Möglichkeiten, dass dieses Ereignis eintritt, z. B. eine Zwei im ersten und eine Fünf im zweiten Wurf. $A \cup B$ ist das Ereignis, dass man im ersten Wurf eine gerade Augenzahl wirft oder im zweiten eine größere Augenzahl als 3. Diese Menge hat 27 Elemente. A^c hat wie A 18 Elemente. Es handelt sich um das Ereignis, dass die erste Augenzahl *nicht* gerade (also ungerade) ist. A und B sind natürlich nicht disjunkt, wir haben $A \cap B$ ja bereits ermittelt. Die Ereignisse schließen sich nicht aus, denn wenn man etwa erst eine Sechs und dann eine Vier würfelt, dann sind sowohl A als auch B eingetreten. Ferner ist weder A eine Teilmenge von B noch umgekehrt. Das liegt auch nahe, weil wir davon ausgehen, dass die Augenzahl des zweiten Wurfes nicht von der des ersten abhängt und die des ersten natürlich auch nicht von der des zweiten.

Lösung 37.3. $\emptyset$ steht für etwas, das garantiert nicht passieren kann, während Ω ein Ereignis ist, das mit Sicherheit eintreten wird. Man spricht auch vom *unmöglichen* bzw. vom *sicheren* Ereignis.

Lösung 37.4. Die Wahrscheinlichkeiten sind

$$P(A) = P(B) = P(A^c) = 18/36 = 1/2 = 50\%,$$

$$P(A \cap B) = 9/36 = 1/4 = 25\% \quad \text{und}$$

$$P(A \cup B) = 27/36 = 3/4 = 75\%.$$

Lösung 37.5. Das sollte so aussehen:

$$A = \{2, 4, 6\}$$

$$B = \{1, 2\}$$

$$P(A) = P(\{2\}) + P(\{4\}) + P(\{6\}) = 2 \cdot 3/20 + 1/4 = 55\%$$

$$P(B) = P(\{1\}) + P(\{2\}) = 2 \cdot 3/20 = 30\%.$$

Bei einem fairen Würfel hätten wir $P(A) = 3/6 = 50\%$ und $P(B) = 2/6 \approx 33\%$.

Lösung 37.6. Wenn man es ausführlich aufschreibt, so erhält man:

$$\bigcap_{n=2}^4 A_{2n} = A_4 \cap A_6 \cap A_8$$

Es geht also um die Teiler, die die drei Zahlen 4, 6 und 8 gemeinsam haben. Das sind 1 und 2, also wäre $\{1, 2\}$ die korrekte Antwort gewesen.

Lösung 37.7. In diesem Durchschnitt sind die Zahlen enthalten, die Teiler von *jeder* positiven natürlichen Zahl n sind. Die einzige Zahl mit dieser Eigenschaft ist die Eins, also ist die Antwort $\{1\}$.

Lösung 37.8. Zunächst zur Schreibweise: Für $\bigcup \mathcal{A}$ kann man auch $\bigcup_{n=1}^{\infty} A_n$ schreiben, für $\bigcap \mathcal{A}$ auch $\bigcap_{n=1}^{\infty} A_n$. Ersetzt man $\mathbb{N}^+$ durch $\mathbb{N}$, so wird daraus $\bigcup_{n=0}^{\infty} A_n$ bzw. $\bigcap_{n=0}^{\infty} A_n$.

Jede positive natürliche Zahl n ist in mindestens einer³¹ der Mengen von $\mathcal{A}$ enthalten, nämlich in $[n]$. Also haben wir offenbar $\bigcup \mathcal{A} = \mathbb{N}^+$. Die einzige Zahl, die in *allen* Mengen von $\mathcal{A}$ enthalten ist, ist hingegen die Eins, also ergibt sich $\bigcap \mathcal{A} = \{1\}$.

Mit $\mathcal{B} = \{[n] : n \in \mathbb{N}\}$ kommt die leere Menge hinzu (siehe Aufgabe 4.6), d. h., wir haben $\mathcal{B} = \mathcal{A} \cup \{\emptyset\}$. Für die Vereinigungsmenge ändert das Hinzufügen der leeren Menge gar nichts, $\bigcup \mathcal{B} = \bigcup \mathcal{A} = \mathbb{N}^+$, aber die Schnittmenge ändert sich, weil es nun *keine* Zahl gibt, die in allen Mengen von $\mathcal{B}$ enthalten ist: $\bigcap \mathcal{B} = \emptyset$.

Lösung 37.9. Hier ergibt sich $\bigcup \mathcal{A} = \mathbb{N}$ und $\bigcap \mathcal{A} = \emptyset$.

Lösung 37.10. Das ist natürlich die leere Menge.

Lösung 37.11. $A_1 = A_2 = [3, 5]$, $A_3 = [4, 5]$, $A_4 = A_6 = \mathbb{N}$, $A_5 = \{0\}$, $A_7 = [0, \infty)$, $A_8 = A_7 \setminus \mathbb{N}$, $A_9 = A_{10} = A_{11} = \emptyset$.

Lösung 37.12. So sollte es aussehen:

$$A \setminus \bigcap_{k=1}^n B_k = \bigcup_{k=1}^n (A \setminus B_k)$$

$$A \setminus \bigcup_{k=1}^n B_k = \bigcap_{k=1}^n (A \setminus B_k)$$

Die Begründung für die Korrektheit folgt im Anschluss an diese Aufgabe.

Lösung 37.13. Im ersten Fall wird eine Menge von Paaren definiert und keine Menge von Zahlen. Im zweiten Fall wird eine Menge von Mengen definiert. Man kann die zweite Variante aber reparieren, indem man $\bigcup \{ \{n, -n\} : n \in \mathbb{N} \}$ schreibt. Das ist dann tatsächlich die Menge der ganzen Zahlen.

³¹Sogar in unendlich vielen, aber eine reicht.

Lösung 37.14. Konnten Sie die Aufgabe lösen? Es ist keine Schande, wenn man so eine Aufgabe nicht lösen kann. Aber wenn Sie etwas lernen wollen, dann ist es dumm, wenn Sie es nicht zumindest probiert haben. Wie viele mathematische Aufgaben ist diese eine, bei der es keinen vorgefertigten Lösungsweg gibt, den man nur nachäffen muss. Häufig hilft nur Nachdenken und – das war ja auch der Tipp – Herumprobieren. Man kommt nicht umhin, einen Stift in die Hand zu nehmen und herumzukritzeln. Was weiß ich? Was fällt mir ein? Stimmt das auch?

Hier kommt man auf diesem Weg schnell zum Ziel. Wegen der ersten Eigenschaft, ist die kleinste Menge, die infrage kommt, $\mathcal{A} = \{\mathbb{R}\}$. Und die tut es auch schon!

Lösung 37.15. Das Ereignis kann man durch die Menge

$$A = \{2n : n \in \mathbb{N}^+\} = \bigcup_{n=1}^{\infty} \{2n\}$$

darstellen. Rechnet man so wie eben, so erhält man als Wahrscheinlichkeit

$$P(A) = \sum_{n=1}^{\infty} P(\{2n\}) = \sum_{n=1}^{\infty} 1/2^{2n} = \sum_{n=1}^{\infty} 1/4^n = \frac{1}{3}.$$

Wenn man so vorgeht, berechnet man in gewissem Sinne Grenzwerte von Wahrscheinlichkeiten. In diesem Fall handelt es sich um den Grenzwert der Folge

$$P(\{2\}), P(\{2, 4\}), P(\{2, 4, 8\}), P(\{2, 4, 8, 10\}), \dots$$

Lösung 37.17. In der Definition steht, dass eine Menge A dann zu $\mathcal{A}$ gehört, wenn ihr Durchschnitt mit $\{1, 2\}$ zwei Elemente hat oder keins; mit anderen Worten: wenn entweder 1 und 2 beide zu A gehören oder beide nicht.

Lösung 37.19. Nur im dritten Fall handelt es sich um eine σ -Algebra, wie man leicht überprüfen kann. Im ersten und vierten Fall gehört $\emptyset$ nicht zu $\mathcal{A}$. Im zweiten Fall sind z. B. $\{1, 2\}$ und $\{2, 3\}$ Elemente von $\mathcal{A}$, ihre Vereinigung ist es aber nicht. Im fünften Fall sind die Singletons $\{0\}$, $\{2\}$, $\{4\}$, $\{6\}$ und so weiter alle Elemente von $\mathcal{A}$, ihre Vereinigung ist jedoch die Menge aller geraden natürlichen Zahlen, die nicht zu $\mathcal{A}$ gehört.

Lösung 37.20. Man könnte der Einfachheit halber $\Omega = [52]$ setzen. Ordnet man den Assen die Zahlen 1 bis 4 zu, so ergibt sich die gesuchte Wahrscheinlichkeit als

$$P(\{1, 2, 3, 4\}) = \frac{|\{1, 2, 3, 4\}|}{|\Omega|} = \frac{4}{52} = \frac{1}{13} \approx 7.7\%$$

Lösung 37.21. In diesem Fall bietet es sich an,

$$\Omega = \{X \subseteq [52] : |X| = 2\}$$

zu setzen. Aus der Schule sollte bekannt sein, dass diese Menge $\binom{52}{2} = 1326$ Elemente hat, wobei mit $\binom{52}{2}$ der **Binomialkoeffizient** gemeint ist. Bezeichnet man dieASSE wie in Aufgabe 37.20, so sieht das uns interessierende Ereignis so aus:

$$A = \{X \subseteq \{1, 2, 3, 4\} : |X| = 2\} = \{\{1, 2\}, \{1, 3\}, \{1, 4\}, \{2, 3\}, \{2, 4\}, \{3, 4\}\}$$

A hat sechs Elemente. $6/1326 = 1/221 \approx 0.5\%$ ist also die gesuchte Wahrscheinlichkeit für zwei Asse.

Beachten Sie, dass in diesem Wahrscheinlichkeitsraum die Ergebnisse selbst Mengen von Zahlen sind, so dass die Ereignisse Mengen von solchen Mengen sind. Das mag etwas verwirrend sein, wenn man es zum ersten Mal sieht.

Man hätte es aber auch anders machen und mit

$$\Omega^* = \{(x_1, x_2) \in [52]^2 : x_1 \neq x_2\}$$

anfangen können. Ein Paar (x_1, x_2) steht dabei dafür, dass x_1 die erste gezogene Karte ist und x_2 die zweite. Ω^* hat die Mächtigkeit $52 \cdot 51$, weil es hier um **Variationen ohne Wiederholungen** geht. Das für uns relevante Ereignis würde in diesem Fall aber anders aussehen:

$$A^* = \{(1, 2), (1, 3), (1, 4), (2, 3), (2, 4), (3, 4), (2, 1), (3, 1), (4, 1), (3, 2), (4, 2), (4, 3)\}$$

Als Wahrscheinlichkeit ergibt sich $|A^*|/|\Omega^*| = 12/(52 \cdot 51)$ und das ist natürlich auch $1/221$.

Lösung 37.22. Die beiden Ereignisse sind nicht disjunkt, daher kann man die beiden Wahrscheinlichkeiten nicht einfach addieren. Zum Beispiel ist die Wahrscheinlichkeit, *kein* Bild zu ziehen, $1 - 3/13 = 10/13$ und die Wahrscheinlichkeit, *nicht* Herz zu ziehen, beträgt $1 - 1/4 = 3/4$. Wenn man das einfach addieren könnte, dann wäre die Wahrscheinlichkeit für eine Karte, die kein Bild *oder* nicht Herz ist, $10/13 + 3/4 = 53/52$, also größer als eins! Was sollte das bedeuten? Und wäre dann die Wahrscheinlichkeit, eine Karte wie Herz Bube zu ziehen, negativ? Das ist natürlich Unsinn.

Die richtige Antwort auf die Frage liefert Punkt (v) von Lemma 37.1. Ist A das Ereignis „Bild“ und B das Ereignis „Herz“, so wollen wir $P(A \cup B)$ berechnen und müssen daher zusätzlich zu $P(A)$ und $P(B)$ noch $P(A \cap B)$ kennen. Bei $A \cap B$ geht es um Karten, die sowohl die Farbe Herz haben als auch Bilder sind. Davon gibt es drei. Also ergibt sich

$$\begin{aligned} P(A \cup B) &= P(A) + P(B) - P(A \cap B) \\ &= 12/52 + 13/52 - 3/52 = 22/52 \approx 42\%. \end{aligned}$$

Lösung 37.23. Wir modellieren das als Laplace-Experiment mit $\Omega = [6]^3$. Man könnte sich nun mühsam überlegen, wie viele verschiedene Möglichkeiten es gibt, eine bestimmte Augenzahl k ab fünf zu erreichen, und diese Möglichkeiten dann alle addieren. Zum Beispiel gibt es für fünf Augen die $M_5 = 6$ Möglichkeiten $(1, 1, 3)$, $(1, 3, 1)$, $(3, 1, 1)$, $(1, 2, 2)$, $(2, 1, 2)$ und $(2, 2, 1)$. Das ergäbe diese Tabelle:

k	5	6	7	8	9	10	11	12	...	17	18
M_k	6	10	15	21	25	27	27	25	...	3	1

Als Summe ergibt sich 212 und damit als Wahrscheinlichkeit $212/6^3 \approx 98\%$.

Ein häufig sinnvoller „Trick“ ist jedoch, sich zu überlegen, ob die „umgekehrte“ Wahrscheinlichkeit nicht einfacher zu berechnen ist. Wenn die Gesamtaugenzahl *nicht* mindestens fünf ist, dann kann sie nur drei oder vier sein. Dafür gibt es vier Möglichkeiten, nämlich $(1, 1, 1)$, $(1, 1, 2)$, $(1, 2, 1)$ und $(2, 1, 1)$. Also ergibt sich die gesuchte Wahrscheinlichkeit als $1 - 4/6^3$. Das ist natürlich viel einfacher.

Lösung 37.24. Das Modell ist richtig. Man kann sich das dadurch klarmachen, dass man mit einem Würfel drei Würfe macht statt mit drei Würfeln einen. Oder man stellt sich vor, die drei Würfel hätten unterschiedliche Farben. Dann wäre der eine Wurf $(1, 4, 5)$ und der andere $(1, 5, 4)$.

Lösung 37.25. Betrachtet man die beiden Alternativen (i) und (ii) als Ereignisse A und B , so gilt offensichtlich $B \subseteq A$ und damit $P(B) \leq P(A)$ nach (iv) von Lemma 37.1. Die zweite Alternative kann also unmöglich *wahrscheinlicher* sein als die erste. (Wenn Sie daran interessiert sind, wie Kahneman diesen Trugschluss erklärt, informieren Sie sich über *Repräsentativitätsheuristik*.)

Lösung 37.26. Wir setzen natürlich $\Omega = [6]^2$ und betrachten das entsprechende Laplace-Experiment. A sei das Ereignis „Augensumme ist prim“ und B das Ereignis, dass der erste Wurf eine Zwei ist. Offenbar gilt $B = \{2\} \times [6]$, also $|B| = 6$, und $A \cap B = \{(2, 1), (2, 3), (2, 5)\}$, also $|A \cap B| = 3$. Damit ergibt sich $P(A|B) = 3/6 = 50\%$.

Übrigens liegt $P(A)$ bei etwa 42 %. Vielleicht überprüfen Sie das zur Übung mal.

Lösung 37.27. Wir haben es wieder mit dem bekannten Laplace-Experiment zu tun. Ich notiere ohne weitere Kommentare einfach die notwendigen Schritte.

$$\begin{aligned}\Omega &= [6]^2 \\ |\Omega| &= 6^2 = 36 \\ A &= \{(x_1, x_2) \in \Omega : x_1 \neq x_2\} = \Omega - \{(x_1, x_1) : x_1 \in [6]\} \\ |A| &= 36 - 6 = 30 \\ B &= \{(x_1, x_2) \in \Omega : x_1 + x_2 = 6\} \\ &= \{(1, 5), (2, 4), (3, 3), (4, 2), (5, 1)\} \\ |B| &= 5 \\ A \cap B &= \{(1, 5), (2, 4), (4, 2), (5, 1)\} \\ |A \cap B| &= 4\end{aligned}$$

Also ergeben sich $P(A|B) = 4/5 = 80\%$ und $P(B|A) = 4/30 \approx 13\%$. Natürlich sind $P(A|B)$ und $P(B|A)$ im Allgemeinen nicht identisch.

Lösung 37.28. Die folgende Übersicht zeigt alle 16 Möglichkeiten, wobei 0 für *Zahl* steht und 1 für *Kopf*. Die Ergebnisse, die zu A gehören, sind blau umrandet, die zu B gehörigen braun hinterlegt.

0000	0001	0010	0011
0100	0101	0110	0111
1000	1001	1010	1011
1100	1101	1110	1111

Man erhält also:

$$\begin{aligned}
 P(A) &= 4/16 = 1/4 = 25\% & P(B) &= 5/16 \approx 31\% \\
 P(A|B) &= 1/5 = 20\% & P(B|A) &= 1/4 = 25\%
 \end{aligned}$$

Lösung 37.29. Das geht so:

$$\begin{aligned}
 A &= \mathbb{N}^+ \setminus \{1, 2\} = \{3, 4, 5, 6, \dots\} \\
 P(A) &= 1 - P(\{1, 2\}) = 1 - (1/2 + 1/4) = 1/4 \\
 B &= \{1, 4, 7, \dots\} = \{3k + 1 : k \in \mathbb{N}\} \\
 P(B) &= \sum_{k=0}^{\infty} \frac{1}{2^{3k+1}} = \frac{1}{2} \cdot \sum_{k=0}^{\infty} \frac{1}{8^k} = \frac{4}{7} \\
 A \cap B &= \{4, 7, 10, \dots\} = B \setminus \{1\} \\
 P(A \cap B) &= P(B) - P(\{1\}) = 4/7 - 1/2 = 1/14 \\
 P(A|B) &= P(A \cap B) / P(B) = 1/14 \cdot 7/4 = 1/8
 \end{aligned}$$

Die Wahrscheinlichkeit dafür, dass man mindestens drei Würfe braucht, beträgt also 25%. Weiß man jedoch, dass das Ereignis B eintreten wird, so liegt diese Wahrscheinlichkeit bei nur noch 12.5%.

Lösung 37.30. Nein. Beachten Sie, dass es 38 mögliche Ergebnisse (Null, Doppel-null und die Zahlen von 1 bis 36) gibt. Nennen wir die Ereignisse aus der Aufgabenstellung A und B , so gilt $P(A) = P(B) = 12/38$ und $P(A \cap B) = 4/38 \neq P(A) \cdot P(B)$.

Lösung 37.31. Es gilt $P(A) = P(B) = 1/2$ und $P(A \cap B) = 1/4 = P(A)P(B)$. Die Ereignisse sind also unabhängig. Wenn Sie sich darüber wundern, dann liegt es daran, dass Sie *kausale* Abhängigkeit und *stochastische* Abhängigkeit durcheinanderbringen. Die Gesamtaugenanzahl hängt natürlich vom Ergebnis des ersten Wurfs (und auch von dem des zweiten) ab. Aber das Wissen darüber, ob das Ereignis A eingetreten ist, liefert Ihnen keine Informationen über die Wahrscheinlichkeit des Eintretens von B , die Sie nicht vorher schon hatten.

Lösung 37.32. Die drei Ereignisse haben alle die Wahrscheinlichkeit $1/2$, denn es gilt $|A| = |B| = |C| = 18$. Außerdem haben wir:

$$\begin{aligned}
 A \cap B &= \{1, 3, 5, 7, 9, 12, 14, 16, 18\} \\
 A \cap C &= \{2, 4, 6, 8, 10, 12, 14, 16, 18\} \\
 B \cap C &= \{12, 14, 16, 18, 30, 32, 34, 36\} \\
 |A \cap B| &= |A \cap C| = 9 \\
 |B \cap C| &= 8
 \end{aligned}$$

$$P(A)P(B) = P(A)P(C) = P(B)P(C) = 1/4$$

$$P(A \cap B) = P(A \cap C) = 1/4$$

$$P(B \cap C) = 2/9$$

Stochastische Unabhängigkeit ist also nicht **transitiv**: Obwohl B und A sowie A und C jeweils unabhängig sind, sind B und C *nicht* unabhängig.

Lösung 37.33. Offenbar gilt jeweils $P(A_i) = 1/2$ und man kann sich leicht überzeugen, dass die drei Ereignisse **paarweise** unabhängig sind. (Tun Sie das!) Aber $A_1 \cap A_2 \cap A_3$ ist das **unmögliche Ereignis** $\emptyset$, während $P(A_1)P(A_2)P(A_3)$ nicht null ist. Aus paarweiser Unabhängigkeit folgt also *nicht* die vollständige Unabhängigkeit!

Lösung 38.1. Am einfachsten ist es wohl, $\Omega = \{0, 1\}^3$ zu setzen, wobei die Eins für *Bild* steht. Dann kann man X durch $(x_1, x_2, x_3) \mapsto x_1 + x_2 + x_3$ definieren. Der Wertebereich von X (also die Menge der Realisierungen) ist $\{0, 1, 2, 3\}$.

Lösung 38.2. Das Ereignis ist $\{(0, 0, 0), (1, 0, 0), (0, 1, 0), (0, 0, 1)\}$ und hat vier Elemente. Da es sich um ein Laplace-Experiment handelt, ist die gesuchte Wahrscheinlichkeit $4/8 = 1/2$.

Lösung 38.3. Sinnvollerweise sollte $\Omega = \{0, 1, 2, 3, \dots, 36\}$ die Ergebnismenge sein. Die Zufallsvariable für *Manque* sieht dann so aus:

$$X_M : \begin{cases} \Omega \rightarrow \mathbb{R} \\ \omega \mapsto \begin{cases} 5 & 1 \leq x \leq 18 \\ -5 & x = 0 \text{ oder } x > 18 \end{cases} \end{cases}$$

Und für *Les trois premiers* so:

$$X_L : \begin{cases} \Omega \rightarrow \mathbb{R} \\ \omega \mapsto \begin{cases} 55 & x \leq 2 \\ -5 & x > 2 \end{cases} \end{cases}$$

Lösung 38.4. Damit der Wert von X größer als 10 ist, muss eines der Ereignisse $(6, 6)$, $(5, 6)$ oder $(6, 5)$ eintreten. Es ergibt sich:

$$\begin{aligned} P(X \geq 11) &= P(X = 11) + P(X = 12) = P(\{(5, 6), (6, 5)\}) + P(\{(6, 6)\}) \\ &= 2 \cdot \frac{3}{20} \cdot \frac{1}{4} + \frac{1}{4} \cdot \frac{1}{4} = \frac{11}{80} = 13.75\% \end{aligned}$$

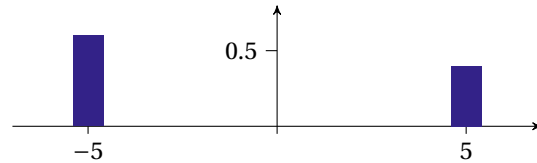
Bei einem fairen Würfel wären das nur etwa 8.3% gewesen.

Lösung 38.5. Mit $I = [6]$ und $x_i = i$ sowie $p_i = 1/6$ für $i \in I$ ergibt sich der Erwartungswert $E(X) = 7/2 = 3.5$. Aber es wird natürlich niemals 3.5 gewürfelt werden. Genauso wenig wird es jemals ein Roulettespiel geben, bei dem das Casino den siebenunddreißigsten Teil eines Euros einnimmt. Der Erwartungswert von X ist trotz des Namens im Allgemeinen nicht ein zu erwartender *Wert*, den X tatsächlich annehmen wird, sondern es handelt sich um so etwas wie „den langfristig zu erwartenden Mittelwert“.

Lösung 38.6. Hier haben wir $p_6 = 1/4$ und $p_i = 3/20$ für $i \in [5]$. Damit ergibt sich als Erwartungswert

$$\left(\sum_{i=1}^5 \frac{3}{20} \cdot i\right) + \frac{1}{4} \cdot 6 = \frac{15}{4} = 3.75.$$

Lösung 38.7. Der Wertebereich von X ist $\{-5, 5\}$, weil der Spieler entweder fünf Euro verliert oder fünf Euro gewinnt. Mit $x_1 = -5$ und $x_2 = 5$ erhalten wir die Wahrscheinlichkeiten $p_1 = P(X = x_1) = 3/5$ und $p_2 = 2/5$.



Der Erwartungswert ist $E(X) = 3/5 \cdot (-5) + 2/5 \cdot 5 = -1$. Im Schnitt verliert der durch die gefälschte Münze getäuschte Spieler somit einen Euro pro Münzwurf.

Lösung 38.8. Mit $\Omega = [6]^2$ und Laplace-Wahrscheinlichkeit würde man X als

$$X: \begin{cases} \Omega \rightarrow \mathbb{R} \\ (a, b) \mapsto \begin{cases} 20 & 4 \mid (a+b) \\ -(a+b) & \text{sonst} \end{cases} \end{cases}$$

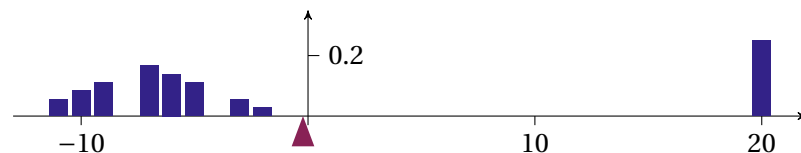
definieren. Damit hätte man z. B. $A = \{\omega \in \Omega : X(\omega) = -3\} = \{(1, 2), (2, 1)\}$ und $P(X = -3) = P(A) = 2/36 = 1/18$.

Insgesamt ergeben sich die folgenden Werte:

i	1	2	3	4	5	6	7	8	9
x_i	20	-2	-3	-5	-6	-7	-9	-10	-11
p_i	1/4	1/36	1/18	1/9	5/36	1/6	1/9	1/12	1/18
$p_i x_i$	5	-1/18	-1/6	-5/9	-5/6	-7/6	-1	-5/6	-11/18

Den Erwartungswert von X erhält man nun durch Summieren über die letzte Zeile. Das Ergebnis ist $-2/9$, was einem durchschnittlichen Verlust von über 20 Cent pro Spiel entspricht. Sie sollten also lieber nicht spielen!

Instruktiv ist hier übrigens ein Blick auf das Stabdiagramm mit eingezeichnetem Erwartungswert (an der Spitze des roten Dreiecks):



Wenn Sie sich die waagerechte Achse als eine Wippe vorstellen und den blauen Stäben ein Gewicht proportional zu ihrer Höhe zuordnen, dann sitzt der Erwartungswert gerade so, dass die Konstruktion ausbalanciert ist.

Lösung 38.9. Der Wertebereich von X ist $\{10, 15, 25, 30\}$. Für 10 und 30 gibt es jeweils nur eine Möglichkeit (beide Fünfeuroscheine bzw. die beiden „großen“ Scheine), während es für 15 und 25 je zwei gibt (weil es zwei Fünfeuroscheine gibt). Die Wahrscheinlichkeiten sehen daher so aus:

i	1	2	3	4
x_i	10	15	25	30
p_i	1/6	1/3	1/3	1/6
$p_i x_i$	5/3	5	25/3	5

Als Summe über die letzte Zeile ergibt sich $E(X) = 20$.

Lösung 38.10. In PYTHON könnte man es folgendermaßen machen, wobei n die Anzahl der Durchgänge ist:

```
from random import randint

def test(n):
    c = n
    for _ in range(n):
        while randint(0,1) != 1:
            c += 1
    return c/n
```

Lösung 38.11. Nach Satz 32.2 gilt $\sum_{k=0}^m 2^k = 2^{m+1} - 1$. Damit erhält man erst

$$\begin{aligned} \sum_{k=1}^n k 2^{n-k} &\stackrel{(*)}{=} \sum_{m=0}^{n-1} \sum_{k=0}^m 2^k = \sum_{m=0}^{n-1} (2^{m+1} - 1) = 2 \left(\sum_{m=0}^{n-1} 2^m \right) - n \\ &= 2(2^n - 1) - n = 2^{n+1} - n - 2, \text{ dann} \\ \sum_{k=1}^n \frac{k}{2^k} &= \frac{\sum_{k=1}^n k 2^{n-k}}{2^n} = \frac{2^{n+1} - n - 2}{2^n} = 2 - \frac{n+2}{2^n} \text{ und schließlich} \\ \sum_{k=1}^{\infty} \frac{k}{2^k} &= \lim_{n \rightarrow \infty} \sum_{k=1}^n \frac{k}{2^k} = \lim_{n \rightarrow \infty} \left(2 - \frac{n+2}{2^n} \right) = 2. \end{aligned}$$

Die mit einem Stern markierte Umformung ergibt sich durch Umsortieren der Summanden. Für $n = 3$ sieht das z. B. so aus:

$$\begin{aligned} \sum_{k=1}^3 k 2^{3-k} &= 1 \cdot 2^{3-1} + 2 \cdot 2^{3-2} + 3 \cdot 2^{3-3} \\ &= 2^{3-1} + (2^{3-2} + 2^{3-2}) + (2^{3-3} + 2^{3-3} + 2^{3-3}) \\ &= (2^{3-3} + 2^{3-2} + 2^{3-1}) + (2^{3-3} + 2^{3-2}) + 2^{3-3} \\ &= (2^0 + 2^1 + 2^2) + (2^0 + 2^1) + 2^0 \\ &= \sum_{k=0}^2 2^k + \sum_{k=0}^1 2^k + \sum_{k=0}^0 2^k = \sum_{m=0}^2 \sum_{k=0}^m 2^k \end{aligned}$$

Lösung 38.12. Es ergibt sich:

$$E(aX) = p_1 \cdot ax_1 + p_2 \cdot ax_2 + \dots + p_n \cdot ax_n$$

$$aE(X) = a(p_1 \cdot x_1 + p_2 \cdot x_2 + \cdots + p_n \cdot x_n)$$

Klammert man im ersten Ausdruck a aus, erhält man den zweiten. Und das Ausklammern von konstanten Faktoren ist nach Lemma 32.7 auch bei Reihen erlaubt.

Lösung 38.13. Für die Zufallsvariablen aus der Lösung zu Aufgabe 38.3 ergibt sich $E(X_M) = E(X_L) = -5/37$ und es gilt offenbar $X = X_M + 2X_L$. Damit folgt sofort $E(X) = -5/37 + 2 \cdot (-5/37) = -15/37 \approx -0.4$.

Lösung 38.14. Den Erwartungswert $7/2$ für die Augenzahl eines Würfels kennen wir aus Aufgabe 38.5. Mit Lemma 38.3 folgt daraus $E(X) = 7/2 + 7/2 = 7$.

Y hat die Realisierungen $y_1 = 0$, $y_2 = 1$ und $y_3 = 2$. Wir müssen nun die 36 möglichen Fälle durchgehen und erhalten:

i	1	2	3
y_i	0	1	2
$P(Y = y_i)$	1/4	1/2	1/4
$P(Y = y_i)y_i$	0	1/2	1/2

Der Erwartungswert von Y ist also 1. Damit haben wir $E(X)E(Y) = E(XY) = 7$. Für $Z = XY$ wird es ein bisschen aufwendiger:

i	1	2	3	4	5	6	7	8	9	10	11
z_i	0	3	5	7	8	9	11	12	16	20	24
$P(Z = z_i)$	1/4	1/18	1/9	1/6	1/36	1/9	1/18	1/18	1/12	1/18	1/36
$P(Z = z_i)z_i$	0	1/6	5/9	7/6	2/9	1	11/18	2/3	4/3	10/9	2/3

Als Summe über die letzte Zeile ergibt sich $E(XY) = 15/2$.

Lösung 38.15. $X_1 = 1$ kann man als $X_1 \in \{1\}$ schreiben, $X < 4$ als $X \in \{2, 3\}$. $X_1 = 1$ und $X < 4$ haben die Wahrscheinlichkeiten $1/2$ und $1/12$. Der Durchschnitt der beiden Ereignisse ist im üblichen Modell $\{(1, 1), (1, 2)\}$ mit der Wahrscheinlichkeit $1/18$. Und das ist natürlich nicht dasselbe wie das Produkt von $1/2$ und $1/12$. Es ist auch intuitiv klar, dass die Wahrscheinlichkeit für eine kleine Gesamtaugenanzahl sich ändert, wenn man weiß, dass der erste Wurf eine Eins war.

Lösung 38.16. Sei X_i die Augenzahl des i -ten Würfelwurfs. Wir wissen bereits aus Aufgabe 38.5 dass, $E(X_1) = E(X_2) = 7/2$ gilt. Weil X_1 und X_2 unabhängig sind und $Y = X_1 X_2$ gilt, können wir einfach ausrechnen: $E(Y) = E(X_1)E(X_2) = 49/4$.

Die direkte Berechnung führe ich nur exemplarisch vor. Es gibt beispielsweise **drei** Möglichkeiten – nämlich $(1, 4)$, $(4, 1)$ und $(2, 2)$ –, das Produkt 4 zu erhalten. Also gilt $P(Y = 4) = 3/36 = 1/12$. Insgesamt erhält man

$$\frac{(6+12) \cdot 4 + 4 \cdot 3 + (2+3+5+8+10+15+18+20+24+30) \cdot 2 + (1+9+16+25+36)}{36} = \frac{441}{36} = \frac{49}{4}.$$

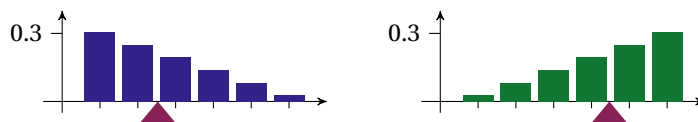
Lösung 38.17. Durch $\omega \mapsto E(X)$ wird eine konstante Funktion definiert, die nach Definition aber auch eine Zufallsvariable ist. Ihr Erwartungswert ist offensichtlich $E(X)$. Nach den Rechenregeln für Erwartungswerte (siehe Lemma 38.3) folgt daher $E(X - E(X)) = 0$.

Lösung 38.18. Mit $\mu_X = 20$ ergänzen wir die Tabelle aus Aufgabe 38.9:

i	1	2	3	4
x_i	10	15	25	30
p_i	1/6	1/3	1/3	1/6
$(x_i - \mu_X)^2$	100	25	25	100
$p_i(x_i - \mu_X)^2$	50/3	25/3	25/3	50/3

Damit ergibt sich $\text{Var}(X) = 50$ als Summe über die letzte Reihe und $\sigma_X = \sqrt{50} \approx 7.1$.

Lösung 38.19. Wie üblich kann man das mit Tabellen berechnen, was ich hier jetzt nicht noch mal vorführe. Das ergibt $\mathbf{E}(X_m) = 91/36 \approx 2.53$, $\mathbf{E}(X_M) = 161/36 \approx 4.47$ und $\text{Var}(X_m) = \text{Var}(X_M) = 2555/1296 \approx 1.97$. Dass X_m und X_M dieselbe Varianz haben, ist nicht weiter erstaunlich, wenn man sich die beiden Stabdiagramme anschaut, die Spiegelbilder voneinander sind.



Die Summe von X_m und X_M ist natürlich einfach die Summe der Augenzahlen, also die Zufallsvariable, die wir in den vorherigen Beispielen X genannt haben. Wie wir bereits wissen (Lemma 38.3), gilt:

$$\mathbf{E}(X_m + X_M) = \mathbf{E}(X) = 7 = 91/36 + 161/36 = \mathbf{E}(X_m) + \mathbf{E}(X_M)$$

Allerdings gilt eine analoge Gleichheit *nicht* für die Varianz:

$$\text{Var}(X_m + X_M) = \text{Var}(X) = 35/6 \neq 2555/648 = \text{Var}(X_m) + \text{Var}(X_M)$$

Lösung 38.20. X_i sei die Bilanz des i -ten Spiels, also $X = X_1 + X_2 + X_3$. X_i kann offenbar die Werte -1 und 1 mit den Wahrscheinlichkeiten $19/37$ bzw. $18/37$ annehmen. Daraus folgt

$$\mathbf{E}(X_i) = -1 \cdot 19/37 + 1 \cdot 18/37 = -1/37 \quad \text{und}$$

$$\text{Var}(X_i) = (-1 + 1/37)^2 \cdot 19/37 + (1 + 1/37)^2 \cdot 18/37 = 1368/1369$$

für $i = 1, 2, 3$. Da die drei Spiele unabhängig sind, ergibt sich:

$$\sigma_X = \sqrt{\text{Var}(X)} = \sqrt{3 \cdot \text{Var}(X_1)} = 6/37 \cdot \sqrt{114} \approx 1.73.$$

Lösung 38.21. Für den ersten Teil wissen wir bereits wegen Lemma 38.3, dass $\mu_{aX} = a\mu_X$ gilt. Damit berechnen wir mittels erneuter Anwendung desselben Lemmas

$$\begin{aligned} \text{Var}(aX) &= \mathbf{E}((aX - \mu_{aX})^2) = \mathbf{E}((aX - a\mu_X)^2) = \mathbf{E}(a^2(X - \mu_X)^2) \\ &= a^2 \mathbf{E}((X - \mu_X)^2) = a^2 \text{Var}(X). \end{aligned}$$

Für den zweiten Teil stelle man sich die Konstante a als eine Zufallsvariable vor, die mit Wahrscheinlichkeit 1 den Wert a annimmt. Ihr Erwartungswert ist dann natürlich a und es folgt $\mu_{X+a} = \mu_X + \mu_a = \mu_X + a$. Das ergibt

$$\begin{aligned}\mathbf{Var}(X + a) &= \mathbf{E}((X + a - \mu_{X+a})^2) = \mathbf{E}((X + a - (\mu_X + a))^2) \\ &= \mathbf{E}((X - \mu_X)^2) = \mathbf{Var}(X).\end{aligned}$$

Im dritten Teil hat man

$$\begin{aligned}\mathbf{Var}(X) &= \mathbf{E}((X - \mathbf{E}(X))^2) = \mathbf{E}(X^2 - 2\mathbf{E}(X)X + \mathbf{E}(X)^2) \\ &= \mathbf{E}(X^2) - \mathbf{E}(2\mathbf{E}(X)X) + \mathbf{E}(\mathbf{E}(X)^2) = \mathbf{E}(X^2) - 2\mathbf{E}(X)\mathbf{E}(X) + \mathbf{E}(X)^2 \\ &= \mathbf{E}(X^2) - \mathbf{E}(X)^2,\end{aligned}$$

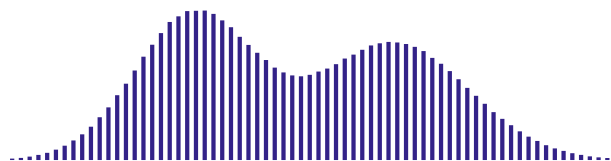
wobei man sich vor Augen halten muss, dass $\mathbf{E}(X)$ hier eine feste Zahl ist.

Für den vierten Teil hat man schließlich unter Ausnutzung des Verschiebungssatzes

$$\begin{aligned}\mathbf{Var}(X + Y) &= \mathbf{E}((X + Y)^2) - \mathbf{E}(X + Y)^2 \\ &= \mathbf{E}(X^2 + 2XY + Y^2) - \mathbf{E}(X + Y)^2 \\ &= \mathbf{E}(X^2) + 2\mathbf{E}(XY) + \mathbf{E}(Y^2) - (\mathbf{E}(X) + \mathbf{E}(Y))^2 \\ &= \mathbf{E}(X^2) + 2\mathbf{E}(X)\mathbf{E}(Y) + \mathbf{E}(Y^2) \\ &\quad - (\mathbf{E}(X)^2 + 2\mathbf{E}(X)\mathbf{E}(Y) + \mathbf{E}(Y)^2) \\ &= \mathbf{E}(X^2) - \mathbf{E}(X)^2 + \mathbf{E}(Y^2) - \mathbf{E}(Y)^2 = \mathbf{Var}(X) + \mathbf{Var}(Y)\end{aligned}$$

Dabei ist die grün hervorgehobene Umformung nur wegen der Unabhängigkeit und Lemma 38.4 erlaubt, während der Rest für beliebige X und Y korrekt ist.

Lösung 38.22. Der Zentrale Grenzwertsatz sagt *nicht*, dass man immer eine Normalverteilung erhält, wenn man nur ausreichend viele Daten hat! Sie erhalten in diesem Fall voraussichtlich eine Verteilung, die ungefähr so aussieht:



Das ist offenbar *keine* Normalverteilung.

Das liegt daran, dass Männer und Frauen im Durchschnitt unterschiedlich groß werden. Mathematisch gesehen haben wir es hier mit einer Folge von Zufallsvariablen zu tun, die *nicht* alle dieselbe Verteilung haben. Bestenfalls haben alle „weiblichen“ Variablen dieselbe Verteilung und alle „männlichen“ ebenfalls dieselbe – aber eben eine andere als die „weiblichen“. Die Voraussetzungen des Zentralen Grenzwertsatzes sind nicht erfüllt und darum ergibt sich nicht die charakteristische Glockenkurve. Und das wird sich auch dann nicht ändern, wenn Sie noch viel mehr Messungen durchführen.

Lösung 39.1. Nein, das ist durchaus üblich. Das Modell setzt nicht eine bestimmte Verteilung voraus, sondern eine ganze Familie von Verteilungen. Für jede Wahl von μ' erhalten wir eine *andere* Normalverteilung. Und der Sinn des Tests ist es, eine Aussage über μ' zu machen. Im Fachjargon ist μ ein *Parameter* der gesuchten Verteilung.

Lösung 39.2. Die Wahrscheinlichkeiten sind ungefähr 95 % und genau 0 %.

Lösung 39.3. Die Standardabweichung ist dann $2.7/\sqrt{18}$ statt $2.7/\sqrt{20}$; die Kurve wird also „breiter“. Der grüne Bereich fängt etwas weiter rechts an, nämlich bei circa 11.45.

Lösung 39.4. Wenn man eine ganze Zahl durch 20 teilt, kommt eine Zahl heraus, die höchstens zwei Nachkommastellen hat. Und wenn es eine zweite Nachkommastelle gibt, die nicht null ist, dann ist es eine Fünf. 11.43 kann nicht das richtige Ergebnis sein. Sollte Ihnen das nicht klar sein, so schauen Sie sich noch einmal Abschnitt 12 an.

Wir rechnen aber für das Beispiel trotzdem mit diesen „krummen“ Werten weiter. Stellen Sie sich vor, es hätte auch „Zwischennoten“ wie 12.4 gegeben. Siehe dazu auch Fußnote 16 auf Seite 307.

Lösung 39.5. 20. Sie müssen dem Programm mitteilen, wie viele „Zufallsexperimente“ Sie durchgeführt haben, weil sich das auf die Standardabweichung des Mittelwerts auswirkt.

Lösung 39.6. Die Nullhypothese ist, dass die Seiten 0 und 1 dieselbe Wahrscheinlichkeit haben, während die Alternativhypothese besagt, dass die Wahrscheinlichkeit für die Seite 1 höher als die für die Seite 0 ist.

Lösung 39.7. Die Nullhypothese ist nach wie vor, dass die Seiten 0 und 1 dieselben Wahrscheinlichkeit haben, während die Alternativhypothese nun besagt, dass die Wahrscheinlichkeiten für die Seiten 0 und 1 unterschiedlich sind. Das ist der bereits angesprochene Unterschied zwischen einem *einseitigen* Test (Aufgabe 39.6) und einem *zweiseitigen*.

Bei einem zweiseitigen Test würden auch Wurfsequenzen mit verdächtig vielen Nullen zur Ablehnung der Nullhypothese führen, während das bei einem einseitigen nicht der Fall wäre. Dafür würden (bei gleichem Signifikanzniveau) bei einem einseitigen Test schon weniger Einsen zur Ablehnung führen als bei einem zweiseitigen.

Lösung 39.8. Wir hatten uns auf Seite 286 bereits überlegt, dass es bei 20 Würfeln 2^{20} verschiedene Abfolgen von Nullen und Einsen gibt, die bei einem fairen Würfel alle dieselbe Wahrscheinlichkeit haben. Eine uns in diesem Fall interessierende Abfolge mit genau fünf Einsen könnte so aussehen:

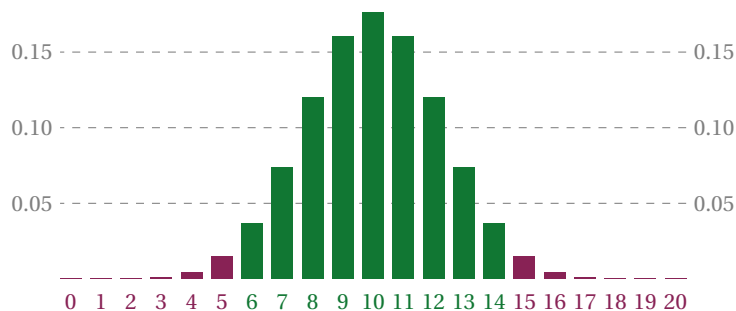
0 0 0 1 1 0 0 0 0 1 0 0 0 0 0 1 0 1 0 0

Die Einsen stehen also an den Positionen 4, 5, 10, 16 und 18. Es gibt $\binom{20}{5}$ Möglichkeiten, fünf Positionen von 20 auszuwählen (siehe Aufgabe 37.21), so dass sich insgesamt die Wahrscheinlichkeit

$$\binom{20}{5} \cdot \frac{1}{2^{20}} = 969/65536 \approx 1.48\%$$

ergibt. Es handelt sich hier um die sogenannte *Binomialverteilung*. Die gesamte Verteilung für 20 Würfe sehen Sie in Aufgabe 39.9.

Lösung 39.9. Die Werte aus der Aufgabenstellung werden in der folgenden Grafik als Säulendiagramm dargestellt.



Die rot gefärbten Ereignisse haben zusammen eine Wahrscheinlichkeit von knapp fünf Prozent. Die Wahrscheinlichkeit, dass eines von denen eintritt, ist also geringer als das Signifikanzniveau. Mit anderen Worten: Wenn bei Ihrer „Studie“ mit 20 Würfeln weniger als sechs oder mehr als 14 Einsen oben lagen, dann müssen Sie (im Rahmen des Hypothesentests) von einer manipulierten Münze ausgehen.

Lösung 39.10. Das Säulendiagramm würde breiter werden, weil es dann mehr Möglichkeiten gäbe. Es gäbe dann allerdings auch mehr Fälle, in denen die Nullhypothese verworfen würde, weil Extremereignisse mit sehr wenigen oder sehr vielen Einsen noch seltener vorkommen würden. Die Wahrscheinlichkeit für einen Fehler 1. Art würde sich *nicht* ändern, denn die wird durch das Signifikanzniveau vorgegeben.

Lösung 39.11. Dafür müsste man für die zu untersuchende Münze die Wahrscheinlichkeiten für Eins und Null kennen. Aber man führt den Test ja gerade deshalb durch, weil man diese *nicht* kennt ...

Literatur

In der folgenden Liste sind diverse Bücher aufgeführt, die alle Mathematik speziell für Studierende der Informatik behandeln. Es ist sicher eine gute Idee, sich ein paar dieser Bücher einmal anzuschauen. Vielleicht verstehen Sie die Erklärungen dort besser als in diesem Skript, vielleicht finden Sie weitere Aufgaben oder vielleicht werden Themen, die hier zu kurz kommen, dort ausführlicher behandelt. Die Links führen jeweils zu den Seiten der Verlage und die Bücher sind danach ausgewählt, dass man Sie als e-Books aus dem Netz der HAW Hamburg (und wohl auch aus den Netzen vieler anderer deutscher Hochschulen) kostenlos herunterladen kann.

- Folkmar Bornemann: [Konkrete Analysis für Studierende der Informatik](#), Springer, 2008.
- Steffen Goebbels / Jochen Rethmann: [Mathematik für Informatiker](#), Springer Vieweg, 2014.
- Peter Hartmann: [Mathematik für Informatiker](#), Springer Vieweg, 2019.
- Stasys Jukna: [Crashkurs Mathematik für Informatiker](#), Teubner, 2008.
- Bernd Kreußler / Gerhard Pfister: [Mathematik für Informatiker](#), Springer, 2009.
- Christoph Meinel / Martin Mundhenk: [Mathematische Grundlagen der Informatik](#), Springer Vieweg, 2015.
- Matthias Schubert: [Mathematik für Informatiker](#), Vieweg + Teubner, 2012.
- Gerald Teschl / Susanne Teschl: [Mathematik für Informatiker Band 1](#) und [Band 2](#), Springer Vieweg, 2013/2014.
- Edmund Weitz: [Konkrete Mathematik \(nicht nur\) für Informatiker](#), Springer Spektrum, 2021.
- Kurt-Ulrich Witt: [Mathematische Grundlagen für die Informatik](#), Springer Vieweg, 2013.

Und im Anschluss noch eine ähnliche Liste. Hier geht es in erster Linie um den Schulstoff, der im Skript vorausgesetzt wird und mit dem viele Studierende erfahrungsgemäß Probleme haben. Wenn Sie Ihre Defizite rechtzeitig (und *nicht* erst kurz vor den Prüfungen!) identifizieren und an ihnen arbeiten wollen, dann kön-

nen die folgenden Bücher dabei sicher helfen. Sie sind nach denselben Kriterien wie oben ausgesucht worden.

- Albrecht Beutelspacher: [Mathe-Basics zum Studienbeginn](#), Springer Spektrum, 2016.
- Erhard Cramer, Johanna Nešlehová: [Vorkurs Mathematik](#), Springer Spektrum, 2015.
- Klaus Dürschnabel et al.: [So viel Mathe muss sein!](#), Springer Spektrum, 2023.
- Arnfried Kemnitz: [Mathematik zum Studienbeginn](#), Springer Spektrum, 2014.
- Marcel Klinger: [Vorkurs Mathematik für Nebenfachstudierende](#), Springer Spektrum, 2015.
- Michael Knorrenschild: [Vorkurs Mathematik](#), Carl Hanser, 2021.
- Werner Poguntke: [Keine Angst vor Mathe](#), Vieweg + Teubner, 2010.
- Sabrina Proß, Thorsten Imkamp: [Brückenkurs Mathematik für den Studieneinstieg](#), Springer Spektrum, 2018.
- Joachim Siegert: [Mathematik trainieren - Brüche, Funktionen und Co.](#), Springer Spektrum, 2018.
- Jürgen Tietze: [Terme, Gleichungen, Ungleichungen](#), Springer Spektrum, 2015.
- Jan van de Craats, Rob Bosch: [Grundwissen Mathematik](#), Springer, 2010.
- Guido Walz, Frank Zeilfelder, Thomas Rießinger: [Brückenkurs Mathematik](#), Springer Spektrum, 2019.



In diesem Zusammenhang empfehle ich auch meine 34-teilige [Videoserie zum Vorkurs Mathematik](#) aus dem Wintersemester 2019.



Index

- Ableitung, 266, 268
 - Ableitung
 - n -te, 268
 - zweite, 268
- Ableitungsfunktion, 268
- Absolutbetrag, 44, 111, 234
- absoluter Fehler, 105
- Abstand, 224
 - euklidischer, 224
- abzählbar, 297
- Additionstheoreme, 160
- Affinkombination, 175
- ähnlich, 145
- Allquantor, 27
- Alternativhypothese, 308
- alternierend, 239
- Ankathete, 148
- Antikommutativität, 385
- äquivalent, 12
- Äquivalenz, 8
- ARCHIMEDES VON SYRAKUS, 271
- Argument, 113, 127
- ARISTOTELES, 7
- Arkuskosinus, 166
- Arkussinus, 166
- Arkustangens, 166
- Assoziativität, 40
- $\operatorname{atan2}$, 168
- atomare Formel, 11
- aufspannen, 213
- aufzählende Mengenschreibweise, 19
- Auslassungspunkte, 32
- Aussage, 7
- Aussageform, 25
- Aussagenvariablen, 11
- BACHMANN, PAUL, 244
- Basis, 46
- Basisvektor, 199
- BAYES, THOMAS, 310
- bayessche Statistik, 310
- bedingte Wahrscheinlichkeit, 291
- Bedingung
 - hinreichende, 10
 - notwendige, 10
- Belegung, 11
- beschränkt, 86, 242
- beschreibende Mengenschreibweise, 33
- bestimmte Divergenz, 243
- Betrag, 44
- Beweis, 6
- BÉZOUT, ÉTIENNE, 61
- Bild, 127, 197
- Bildmenge, 126
- Binomialverteilung, 410
- binäre Exponentiation, 77
- Bisektion, 264
- BOLZANO, BERNARD, 123
- BOOLE, GEORGE, 7, 18
- BOREL, ÉMILE, 288
- borelsche σ -Algebra, 288
- Bruch, 36
- CANTOR, GEORG, 18

- DE MORGAN, AUGUSTUS, *siehe*
Morgan, Augustus De
- de-morgansche Gesetze, 14, 28
- Definition, 6
- Definitionsbereich, 126
- Destrukturierung, 136
- Determinante, 212
- Diagonalmatrix, 204
- Dichtefunktion, 304
- Differentialquotient, 266
- Differenz
mengentheoretische, 21
symmetrische, 23
- differenzierbar, 265
zweimal, 268
- DIRICHLET, PETER GUSTAV LEJEUNE,
123, 273
- Dirichlet-Funktion, 273
- disjunkt, 83
- Disjunktion, 8
- diskrete Zufallsvariable, 297
- Distributivität, 40
- divergent, 238
- divergieren, 243
bestimmt, 243
- Division mit Rest, 55
- Drehrichtung, 139
- Drehspiegelung, 232
- Drehung, 203
- Drei-Finger-Regel, 218
- Dreieck, 142
gleichschenkliges, 143
gleichseitiges, 143
rechtwinkliges, 146
- Dreiecksmatrix, 189
- Dreiecksungleichung, 45, 111, 223,
234
- Durchschnitt, 21
- Ebene, 177
- echt komplex, 109
- echte Teilmenge, 21
- echter Teiler, 57
- echtes Intervall, 97
- Ecke, 142
- Einheit, imaginäre, 107
- Einheitskreis, 350
- Einheitsmatrix, 182
- Einheitsquaternion, 234
- Einheitsvektor, kanonischer, 199
- Element, 18
- elementare
Funktion, 263
Zeilenumformung, 189
- Eliminationsverfahren, 188
- Ereignis, 282, 289
- erfüllbar, 12
- Ergänzungswinkel, 141
- Ergebnis, 281
- Erwartungswert, 298, 304
- erweiterte Koeffizientenmatrix, 187
- EUKLID VON ALEXANDRIA, 59, 137
- euklidische Ebene, 137
- euklidischer
Algorithmus, 59
euklidischer Abstand, 224
- EULER, LEONHARD, 15, 132, 255
- eulersche Zahl, 255
- Existenzquantor, 27
- Exponent, 46
- Exponentialfunktion, 255, 258
- EZ, 189
- FALK, SIGURD, 181
- Falksches Schema, 181
- Fallunterscheidung, 45
- falsch, 7
- fast
alle, 238
unmöglich, 287
- FAULHABER, JOHANNES, 249
- Faulhabersche Formel, 249
- Fehler
1. Art, 311
2. Art, 311
absoluter, 105
relativer, 105
- FERMAT, PIERRE DE, 76
- Fläche
eines Dreiecks, 148
eines Rechtecks, 147
- Folge, 237

geometrische, 241
harmonische, 241
Folnglied, 237
formale Sprache, 11
Formel
 atomare, 11
 der Aussagenlogik, 11
FREGE, GOTTLOB, 25
freie Individuenvariable, 27
Fundamentalsatz
 der Algebra, 116
 der Analysis, 277
 der Arithmetik, 65
Fünfeck, 150
Funktion, 125
 elementare, 263
 gerade, 157
 glatte, 268
 injektive, 128
 konstante, 351
 lineare, 197
 mehrstellige, 130
 periodische, 156
 ungerade, 157
 zweistellige, 130
Funktionsgraph, 127
Funktionswert, 127

ganze Zahl, 35
GAUSS, CARL FRIEDRICH, 15, 68, 107, 131, 188, 305
Gauß-Jordan-Algorithmus, 194
Gauß-Test, 308
Gaußklammer
 obere, 131
 untere, 131
gaußsche
 Summenformel, 248
 Zahlenebene, 110
Gaußsches Eliminationsverfahren, 188
gebundene Individuenvariable, 27
Gegenkathete, 148
Gegenwinkel, 141
genau ein, 31
genau dann, wenn, *siehe* Äquivalenz

Geometrie, 137
geometrische
 Folge, 241
 Reihe, 251
 Summenformel, 248
geordnetes Paar, 119
Gerade, 138
gerade
 Zahlen, 57
gerade Funktion, 157
gerichteter Winkel, 139
glatt, 268
gleichschenkliges Dreieck, 143
gleichseitiges Dreieck, 143
Gleichungssystem, 186
 lineares, 187
GÖDEL, KURT, 15, 67
Gödelisierung, 66, 352
Grad, 140, 262
Grenzwert, 240
groß-Oh, 244
größter gemeinsamer Teiler, 59
größtes Element, 87

Halbgerade, 138
HAMILTON, WILLIAM ROWAN, 232
harmonische
 Reihe, 250
harmonische Folge, 241
Hauptdiagonale, 164, 180
HESSE, OTTO, 229
hessesche Normalenform, 228
Hexagon, 152
HILBERT, DAVID, 15, 25
hinreichende Bedingung, 10
homogen, 188
homogene Koordinaten, 207
Hypotenuse, 146
Hypothesentest, 307

Identität, 127
i. i. d., 305
imaginäre Einheit, 107
Imaginärteil, 107
Implikation, 8
Individuenvariable, 25

- freie, 27
- gebundene, 27
- injektiv, 128
- innere Verknüpfung, 131
- Integral, 273
- Integrand, 274
- Integrationsgrenze, 274
- Integrationsvariable, 274
- integrierbar, 273
- Intervall, 97
 - echtes, 97
- inverse Matrix, 210
- inverses Element, 40, 42
- invertierbar, 210
- irrationale Zahl, 95
- Isometrie, 231
- JORDAN, WILHELM, 194
- Junktor, 8
- KAHNEMAN, DANIEL, 291
- kanonische
 - Primfaktorzerlegung, 65
- kanonischer Einheitsvektor, 199
- kartesisches
 - Koordinatensystem, 121
 - Produkt, 120
- Kathete, 146
- kaufmännisches Runden, 104
- Kehrwert, 42
- Kehrwertregel, 268
- Kettenregel, 269
- Klammerersparnis, 13, 29
- kleiner, 83
- kleinstes
 - gemeinsames Vielfaches, 66
- kleinstes Element, 87
- Koeffizient, 187, 262
- Koeffizientenmatrix, 187
- KOLMOGOROW, ANDREI, 289
- Kolmogorow-Axiome, 289
- Kommutativität, 40
- Komplement, 282
- komplexe Zahl, 107
- Komponente, 119, 180
- Komposition, 132
- kongruent, 143
- Kongruenz, 79
- Kongruenzsystem, 79
- Kongruenzsätze, 143
- konjugiert, 108, 234
- Konjunktion, 8
- konstante Funktion, 351
- konvergent, 238
- konvergieren, 238
- Koordinatensystem, kartesisches, 121
- Korollar, 6
- Kosekans, 149
- Kosinus, 148
- Kotangens, 149
- Kreiszahl, 154
- Kreuzprodukt, 226
- Kryptographie, 55
- LANDAU, EDMUND, 244
- LAPLACE, PIERRE-SIMON, 283
- Laplace-Experiment, 290
- Laplace-Wahrscheinlichkeit, 283
- lateinisches Quadrat, 71
- Laufvariable, 247
- leer, 52
- leere Menge, 19
- leeres Produkt, 65
- leeres Tupel, 119
- LEIBNIZ, GOTTFRIED WILHELM, 123, 272, 274, 277
- Leibnizregel, 267
- Leitkoeffizient, 189
- Lemma, 6
 - von Bézout, 61
- LÉVY, PAUL, 305
- LGS, 187
- Limes, 240
- LINDBERG, JARL WALDEMAR, 305
- lineare Abbildung, 197
- lineares Gleichungssystem, 187
- Linearkombination, 58, 175
- linkseindeutig, 128
- Logarithmus, 257
 - zur Basis a , 258
- Lösung eines Gleichungssystem, 188

- Mächtigkeit, 51
- maplet*, 126
- mathematisches Runden, 104
- Matrix, 179
 - inverse, 210
 - quadratische, 180
- Maximum, 87
- mehrstellige Funktion, 130
- Menge, 18
 - leere, 19
- Mengendiagramm, 20
- Mengenklammern, 19
- Mengenlehre, 18
- Mengenprodukt, 120
- Mengenschreibweise
 - aufzählende, 19
 - beschreibende, 33
- Minimum, 87
- Modell, mathematisches, 281
- MORGAN, AUGUSTUS DE, 14, 28
- Multiplikativität, 45
- n -Tupel, 119
- nach
 - oben beschränkt, 86
 - unten beschränkt, 86
- Nachbarwinkel, 142
- natürliche Zahl, 33
- natürlicher Logarithmus, 257
- Nebenwinkel, 141
- Negation, 8
- negativ, 39, 83
- neutrales Element, 40
- NEWTON, ISAAC, 272, 277
- Newton-Leibniz-Formel, 277
- nichtleer, 88
- nichtnegativ, 39
- Norm, 223
- Normalenform, 227
 - hessesche, 228
- Normalenvektor, 227
- normalverteilt, 305
- normieren, 223, 383
- notwendige Bedingung, 10
- Nullfolge, 240
- Nullfolgenkriterium, 250
- Nullhypothese, 308
- Nullmatrix, 180
- Nullpolynom, 262
- Nullstelle, 264
- Nullvektor, 170
- numerische Quadratur, 278
- o. B. d. A., 66
- obere
 - Dreiecksmatrix, 189
 - Schranke, 86
- Obermenge, 20
- oder, *siehe* Disjunktion
- Ordnung, 244
- ORESME, NIKOLAUS VON, 251
- orientierter Winkel, 139
- orthogonal, 225, 230
- Ortsvektor, 173
- p -Wert, 311
- Paar, geordnetes, 119
- paarweise, 79
- parallel, 138
- Parallelepiped, 212
- Parität, 331
- Partialsumme, 249
- periodische Funktion, 156
- Pfeildiagramm, 125
- Phase, 113
- Pivotelement, 191
- Polarkoordinaten, 159
- Polygon, 150
- Polynom, 262
- positiv, 39, 83
- positiv definit, 222
- Potenz, 45
- Potenzmenge, 50
- prim, 34
- Primfaktor, 64
- Primfaktorzerlegung, 65
- Primteiler, 64
- Primzahl, 34
- Priorität, 13, 29
- Produktregel, 267
- Projektion, 204
- Proposition, 6

- Punkt, 137
 Punkt-Richtungs-Form, 174, 178
 PYTHAGORAS VON SAMOS, 90, 111, 146

 Quadrat, lateinisches, 71
 quadratische Matrix, 180
 Quadratur
 numerische, 278
 Quadratwurzel, 99
 Quadratzahl, 90
 Quadrupel, 119
 Quantor, 27
 Quaternion, 232
 Quersumme, 57
 alternierende, 73
 Quintupel, 119

 Radiant, 155
 rationale Zahl, 36
 Realisierung, 294
 Realteil, 107
 rechter Winkel, 140
 rechtseindeutig, 125
 Rechtssystem, 226
 rechtwinkliges Dreieck, 146
 reelle Zahl, 95
 Regel von Sarrus, 217
 regulär, 210
 Reihe, 249
 geometrische, 251
 harmonische, 250
 Reihenglied, 249
 rein imaginär, 109, 232
 Relation, 125
 relativen Fehler, 105
 Rest, 56
 Restklasse, 70
 Restklassenkörper, 74
 Restklassenring, 71
 Richtungsvektor, 174
 RIEMANN, BERNHARD, 272
 Riemann-Integral, 272
 Runden, 103
 kaufmännisches, 104
 mathematisches, 104

 RUSSELL, BERTRAND, 25

 SARRUS, PIERRE FRÉDÉRIC, 217
 Sarrus, Regel von, 217
 Satz, 6
 des Pythagoras, 146
 des Pythagoras, umgekehrter, 163
 des Thales, 150
 des Theaitetos, 90
 des Euklid, 64
 Säulendiagramm, 297
 Scheitelpunkt, 139
 Scheitelwinkel, 141
 Schenkel, 139
 Scherung, 202
 Schnittmenge, 284
 Schranke, 86
 Sechseck, 152
 Seite, 142
 Sekans, 149
 senkrecht, 225
 σ -Algebra, 287
 signifikant, 310
 signifikante Stellen, 105
 Signifikanzniveau, 309
 Singleton, 19
 singulär, 210
 Sinus, 148
 Sinussatz, 161
 Skalar, 169
 Skalarmultiplikation, 169
 Skalarprodukt, 221
 Skalierung, 202
 Spat, 212
 Spiegelung, 203
 spitzer Winkel, 141
 Sprache, formale, 11
 Stabdiagramm, 297
 Stammfunktion, 276
 Standardabweichung, 301
 Standardnormalverteilung, 305
 Statistik, 306
 bayessche, 310
 statistisch signifikant, 310
 stetig, 261

- stetige Zufallsvariable, 304
- STEVIN, SIMON, 91
- Strahl, 138
- Strahlensatz, 144
- Strecke, 138
- Stufenwinkel, 142
- stumpfer Winkel, 141
- Summe, 249
- Summenformel
 - arithmetische, 248
 - gaußsche, 248
 - geometrische, 248
- Summenschreibweise, 246
- SUN ZI, 79
- symbolische Integration, 278
- symmetrisch, 232

- Tangens, 148
- Tangente, 265
- TARSKI, ALFRED, 25
- Taschenrechner, 155, 330
- Tautologie, 12
- Teilen mit Rest, 55
- Teiler, 57
 - echter, 57
 - größter gemeinsamer, 59
 - trivialer, 57
- Teilmenge, 19
 - echte, 21
- THALES VON MILET, 150
- THEAITETOS, 90
- Theorem, 6
- Transitivität, 58
- Translation, 206
- Transposition, 184
- Tripel, 119
- trivialer Teiler, 57
- Tupel, 119
 - leeres, 119
- Typ, 180

- Umkehrfunktion, 129
- unabhängig, 292, 300
- und, *siehe* Konjunktion
- ungerade
 - Funktion, 157

- Zahlen, 57
- ungerichteter Winkel, 141
- univalent*, 125
- unterbestimmt, 178, 193
- untere Schranke, 86

- Variable, *siehe* Individuenvariable
- Varianz, 301
- Varianzanalyse, 313
- Vektor, 169
- Vektoraddition, 169
- vektoriell, 232
- Vektorprodukt, 226
- Vereinigung, 21
- Vereinigungsmenge, 284
- Verknüpfung, 9
- Verknüpfung, 130
 - innere, 131
- Verschiebungssatz, 302
- verschwinden, 59
- Verteilung, 295
- Verteilungsfunktion, 295
- Vieleck, 150
- Vielfaches, 57
 - kleinstes gemeinsames, 66

- wahr, 7
- Wahrheitstafel, 12
- Wahrheitswert, 7
- Wahrscheinlichkeit, 289
 - bedingte, 291
- Wahrscheinlichkeitsfunktion, 297
- Wahrscheinlichkeitsraum, 289
- Wechselwinkel, 142
- WEIERSTRASS, KARL, 123
- wenn, dann, *siehe* Implikation
- Wert, 249
- Wertebereich, 126
- windschief, 368
- Winkel, 139
 - gerichteter, 139
 - orientierter, 139
 - rechter, 140
 - spitzer, 141
 - stumpfer, 141
 - ungerichteter, 141

- zwischen Vektoren, 224
- wissenschaftliches Runden, 104
- wohldefiniert, 70
- Wohlordnung, 88
- Wurzel, 99
- Zahl
 - echt komplexe, 109
 - ganze, 35
 - gerade, 57
 - irrationale, 95
 - komplexe, 107
 - natürliche, 33
 - rationale, 36
 - reelle, 95
 - rein imaginäre, 109
 - ungerade, 57
- zusammengesetzte, 34
- Zahlengerade, 84
- Zahlentheorie, 55
- ZAPPA, FRANK, 67
- Zeilenstufenform, 189
- Zeilenumformung, 189
- Zentraler Grenzwertsatz, 305
- Zielmenge, 126
- Ziffernblatt, 78
- ZSF, 189
- Zufallsvariable, 294
 - diskrete, 297
 - stetige, 304
- zusammengesetzte Zahl, 34
- zweistellige Funktion, 130
- Zwischenwertsatz, 263

Mathematische Symbole

Es ergibt wenig Sinn, mathematische Symbole alphabetisch zu sortieren. Daher werden sie in der folgenden Liste einfach in der Reihenfolge aufgeführt, in der sie definiert bzw. eingeführt wurden.

$\neg$, 8	a^{-1} , 42
$\wedge$, 8	a/b , 42
$\vee$, 8	$ x $, 44
$\Rightarrow$, 8	x^n , 45
$\Leftrightarrow$, 8	$\mathcal{P}(A)$, 50
$\emptyset$, 19	$ A $, 51
$A \cup B$, 21	2^A , 51
$A \cap B$, 21	$a \bmod m$, 56
$A \setminus B$, 21	$\text{ggT}(a, b)$, 59
$A(x)$, 26	$\text{kgV}(a, b)$, 66
$\forall$, 27	$[a]_m$, 70
$\exists$, 27	$\mathbb{Z}/m\mathbb{Z}$, 71
$A[m]$, 27	$\mathbb{Z}_p$, 74
$\forall x, y \dots$, 29	$a \equiv b \pmod{m}$, 79
$\exists x, y \dots$, 29	$\mathbb{Q}^+$, 83
$a \mid b$, 30	$\mathbb{Q}^-$, 83
$a \nmid b$, 30	$A \dot{\cup} B$, 83
$\exists!$, 31	$a < b$, 83
$\mathbb{N}$, 33	$a \leq b$, 84
$\mathbb{N}^+$, 34	$\max M$, 87
$\mathbb{P}$, 34	$\min M$, 87
$[n]$, 34	$\mathbb{R}$, 95
$\mathbb{Z}$, 35	$\mathbb{R}^+$, 95
$\mathbb{Q}$, 36	$\mathbb{R}^-$, 95
$-a$, 40	$[a, b]$, $[a, b)$, etc., 97
$a - b$, 40	$\sqrt[n]{y}$, 99
$1/a$, 42	$\sqrt{y}$, 99

- $\approx$, 105
 $\operatorname{Re}(z)$, 107
 $\operatorname{Im}(z)$, 107
 $\mathbb{C}$, 107
 i , 107
 $\bar{z}$, 108
 $|z|$, 111
 $\arg(z)$, 113
 $(a_1, a_2, \dots, a_n)$, 119
 $A \times B$, 120
 A^2 , 120
 A^n , 121
 D_f , 126
 W_f , 126
 $f: A \rightarrow B$, 127
 $f(x)$, 127
 $x \mapsto f(x)$, 127
 id_A , 127
 Δ_A , 127
 f^{-1} , 129
 $[x]$, 131
 $\lfloor x \rfloor$, 131
 $R \circ S$, 132
 360° , 140
 $\angle$, 141
 $\sphericalangle$, 141
 $[XY]$, 144
 $|XY|$, 144
 $\sin \alpha$, 148
 $\cos \alpha$, 148
 $\tan \alpha$, 148
 $\cos^3 \alpha$, 148
 π , 154
 $\sin^k x$, 157
 $\arcsin$, 166
 $\arccos$, 166
 $\arctan$, 166
 $\mathbf{a}$, 169
 $\mathbf{0}$, 170
 $\mathbf{p} + \mathbb{R}\mathbf{v}$, 174
 $\mathbb{R}^{m \times n}$, 180
 $\mathbf{0}$, 180
 $\mathbf{E}_n$, 182
 $\mathbb{1}_n$, 182
 $\mathbf{A}^T$, 184
 $\mathbf{v}^T$, 185
 $f[D]$, 197
 $\mathbf{e}_k$, 199
 $[a, b, 1]$, 207
 $\mathbf{A}^{-1}$, 210
 $\det(\mathbf{A})$, 212
 $\mathbf{a} \cdot \mathbf{b}$, 222
 $\mathbf{a}^2$, 222
 $\|\mathbf{a}\|$, 223
 $d(\mathbf{p}, \mathbf{q})$, 224
 $\angle(\mathbf{a}, \mathbf{b})$, 224
 $\mathbf{a} \times \mathbf{b}$, 226
 $\operatorname{Re}(u)$, 232
 $\mathbb{H}$, 232
 $|u|$, 234
 $\bar{u}$, 234
 $(f(n))_{n \in \mathbb{N}}$, 237
 $(f(n))$, 237
 $\lim_{n \rightarrow \infty} a_n$, 240
 $\mathcal{O}(f)$, 244
 $\sum_{k=m}^n a_k$, 246
 $\sum_{k=m}^\infty a_k$, 249
 $n!$, 254
 $\exp(z)$, 255
 e , 255
 e^z , 256
 $\ln$, 257
 $\log$, 257
 a^z , 258
 $\log_a$, 258
 $\lim_{x \rightarrow a} f(x) = L$, 261
 $f + g$, 262
 $\frac{df}{dx}(x_0)$, 265
 $f'(x_0)$, 266
 $f''(x_0)$, 268
 $\frac{d^2 f}{dx^2}(x_0)$, 268
 $A^{\mathbb{C}}$, 282
 $\bigcup \mathcal{A}$, 284
 $\bigcup_{M \in \mathcal{A}} M$, 284
 $\bigcap \mathcal{A}$, 284
 $\bigcap_{M \in \mathcal{A}} M$, 284
 $\bigcup_{n=1}^\infty A_n$, 285
 $\mathcal{B}(\mathbb{R})$, 288
 $P(A|B)$, 291
 $X \leq a$, 295
 $X < a$, 295

$X \in A$, 295

F_X , 295

p_i , 297

$\mathbf{E}(X)$, 298

μ_X , 298

$\mathbf{Var}(X)$, 301

σ_X^2 , 301

σ_X , 301

Φ , 305

$X \sim \mathcal{N}_{\mu, \sigma^2}$, 305