

# Mathevorkurs Technische Fakultät

William Glover   Martin Goth   Thomas Liebetraut  
Patricia Quellmalz   Fabian Wenzelmann   Jan Ole von Hartz



## **Autor\*innen:**

William Glover, Martin Goth, Thomas Liebetraut, Patricia Quellmalz, Fabian Wenzelmann, Jan Ole von Hartz

## **Lektor\*innen:**

Malte Ahl, Christian Schilling, Andreas Steffen, Leonie Feldbusch, Ann-Catherine Nölken, Hans Albert

## **Version:**

2.4.1, September 2025

## **Kontakt:**

Albert-Ludwigs-Universität Freiburg  
Fachschaft TF  
Georges-Köhler-Allee 51  
79110 Freiburg  
`erstiorga@fachschaft.tf.uni-freiburg.de`

## **Lizenz:**

Copyright © 2012 - 2020 William Glover, Martin Goth, Thomas Liebetraut, Patricia Quellmalz, Fabian Wenzelmann, Jan Ole von Hartz.

Dieses Skript ist veröffentlicht unter den Bedingungen der Creative Commons Attribution - ShareAlike 4.0 International (CC BY-SA 4.0) Lizenz.

This work is licensed under a Creative Commons Attribution - ShareAlike 4.0 International (CC BY-SA 4.0) License.

<http://creativecommons.org/licenses/by-sa/4.0/>



# Vorwort

Dieser Vorkurs soll als Einstieg in die Mathematik dienen, die ihr für Euer Studium an der Technischen Fakultät braucht. An der Universität ist die Herangehensweise an Mathe meist anders als in der Schule. In diesem Vorkurs versuchen wir Konzepte, die Ihr teils noch aus der Schule kennt, ähnlich zur Mathevorlesung zu vermitteln. Ziel des Kurses ist es, eine Intuition für wichtige Konzepte zu vermitteln, die Ihr in den ersten Semestern im Studium lernen werdet. Nach dem Vorkurs müsst ihr nicht alles zu 100% reproduzieren können. Ihr solltet aber versuchen, ein Verständnis für die Art und Weise, wie Mathe an der Uni aufgeschrieben wird aufbauen. Hilfreich ist es, eine Intuition für einige Begrifflichkeiten zu entwickeln, damit Ihr in der Vorlesung weniger Schwierigkeiten habt, mitzukommen.

Es ist nicht schlimm, wenn euere Tutorin nicht das gesamte Skript in der Zeit behandeln kann. Ihr könnt das Skript auch nochmal während den Vorlesungen auspacken, weil wir versuchen, Inhalte ausführlicher einzuführen als sie bei manchen der Matheprofessorinnen eingeführt werden.

Viel Spaß beim Mathevorkurs, ihr schafft Mathe I schon.

Hans Albert 2025

Einen Dank an:

William Glover, Martin Goth, Thomas Liebetraut, Patricia Quellmalz, Fabian Wenzelmann für die Erstellung des Mathevorkurs.

Wir danken Desweiteren:

André Blickensdörfer, Layla Franke, Lukas Halbritter, Lukas Klar, Thomas Liebetraut, Patricia Quellmalz, Christian Schilling, Andreas Steffen, Fabian Wenzelmann für die 2. Auflage.

Zu guter letzt einen Dank an:

Ann-Catherine Nölken und Jan Ole von Hartz für die 3. Auflage.



# Inhaltsverzeichnis

|                                                |           |
|------------------------------------------------|-----------|
| <b>1. Definition – Satz – Beweis</b>           | <b>5</b>  |
| <b>2. Mengenlehre</b>                          | <b>9</b>  |
| 2.1. Mengenoperationen . . . . .               | 10        |
| 2.2. Relationen . . . . .                      | 13        |
| 2.2.1. Komposition . . . . .                   | 15        |
| 2.2.2. Eigenschaften von Relationen . . . . .  | 16        |
| <b>3. Zahlenmengen und Operationen</b>         | <b>19</b> |
| 3.1. Zahlenräume . . . . .                     | 19        |
| 3.1.1. Natürliche Zahlen . . . . .             | 19        |
| 3.1.2. Ganze Zahlen . . . . .                  | 21        |
| 3.1.3. Rationale Zahlen . . . . .              | 21        |
| 3.1.4. Reelle Zahlen . . . . .                 | 22        |
| 3.1.5. Komplexe Zahlen . . . . .               | 22        |
| 3.2. Operationen . . . . .                     | 27        |
| 3.2.1. Potenz . . . . .                        | 27        |
| 3.2.2. Wurzel . . . . .                        | 28        |
| 3.2.3. Logarithmus . . . . .                   | 28        |
| <b>4. Analysis</b>                             | <b>31</b> |
| 4.1. Funktionen . . . . .                      | 31        |
| 4.1.1. Eigenschaften von Funktionen . . . . .  | 32        |
| 4.1.2. Transformationen . . . . .              | 34        |
| 4.2. Folgen . . . . .                          | 36        |
| 4.2.1. Konvergenz . . . . .                    | 36        |
| 4.2.2. Wichtige Folgen . . . . .               | 37        |
| 4.2.3. Rechenregeln für Zahlenfolgen . . . . . | 39        |
| 4.2.4. Rekursion . . . . .                     | 39        |
| 4.3. Reihen . . . . .                          | 41        |
| 4.3.1. Summen und Produkte . . . . .           | 41        |
| 4.3.2. Reihen . . . . .                        | 42        |
| 4.3.3. Wichtige Reihen . . . . .               | 43        |
| 4.4. Funktionsklassen . . . . .                | 44        |
| 4.4.1. Lineare Funktionen . . . . .            | 44        |
| 4.4.2. Polynome . . . . .                      | 45        |
| 4.4.3. Wurzelfunktionen . . . . .              | 46        |

|           |                                                           |           |
|-----------|-----------------------------------------------------------|-----------|
| 4.4.4.    | Trigonometrische Funktionen . . . . .                     | 47        |
| 4.4.5.    | Exponentialfunktionen . . . . .                           | 48        |
| 4.4.6.    | Logarithmusfunktionen . . . . .                           | 48        |
| 4.5.      | Nullstellenberechnung . . . . .                           | 49        |
| 4.6.      | Differentialrechnung . . . . .                            | 51        |
| 4.6.1.    | Gängige Ableitungen . . . . .                             | 53        |
| 4.6.2.    | Rechenregeln für die Differentiation . . . . .            | 55        |
| 4.7.      | Integralrechnung . . . . .                                | 58        |
| 4.7.1.    | Integrationsregeln . . . . .                              | 61        |
| <b>5.</b> | <b>Lineare Algebra</b>                                    | <b>65</b> |
| 5.1.      | Vektoren . . . . .                                        | 65        |
| 5.1.1.    | Nomenklatur . . . . .                                     | 65        |
| 5.1.2.    | Vektoroperationen . . . . .                               | 66        |
| 5.1.3.    | Lineare Abhängigkeit . . . . .                            | 69        |
| 5.2.      | Matrizen . . . . .                                        | 70        |
| 5.2.1.    | Nomenklatur . . . . .                                     | 71        |
| 5.2.2.    | Matrizenoperationen . . . . .                             | 71        |
| 5.2.3.    | Berechnung von Determinanten . . . . .                    | 73        |
| 5.3.      | Lineare Gleichungssysteme . . . . .                       | 76        |
| 5.3.1.    | Substitutionsverfahren . . . . .                          | 77        |
| 5.3.2.    | Gaußsches Eliminationsverfahren . . . . .                 | 77        |
| <b>6.</b> | <b>Logik und Beweise</b>                                  | <b>79</b> |
| 6.1.      | Aussagenlogik . . . . .                                   | 79        |
| 6.1.1.    | Junktoren . . . . .                                       | 80        |
| 6.1.2.    | Erfüllbarkeit und Äquivalenzen . . . . .                  | 83        |
| 6.2.      | Logik der ersten Stufe . . . . .                          | 87        |
| 6.3.      | Beweise . . . . .                                         | 88        |
| 6.3.1.    | Wozu Beweisen? . . . . .                                  | 88        |
| 6.3.2.    | Direkter Beweis . . . . .                                 | 88        |
| 6.3.3.    | Beweis durch Widerspruch (Reductio ad absurdum) . . . . . | 89        |
| 6.3.4.    | Vollständige Induktion . . . . .                          | 90        |
| <b>A.</b> | <b>Griechische Buchstaben</b>                             | <b>93</b> |
| <b>B.</b> | <b>Mathematische Symbole</b>                              | <b>95</b> |
| B.1.      | Mengenlehre . . . . .                                     | 95        |
| B.2.      | Analysis und Arithmetik . . . . .                         | 96        |
| B.3.      | Vektorrechnung . . . . .                                  | 97        |
| B.4.      | Logik . . . . .                                           | 97        |

# 1. Definition – Satz – Beweis

Jeder Beweis ist die Zurückführung des Zweifelhaften auf ein Anerkanntes.

---

Arthur Schopenhauer

In der Schule wird einem hauptsächlich beigebracht, *wie* man etwas berechnet. Zum Beispiel sollte jeder\*jedem bekannt sein, dass es den Satz des Pythagoras gibt und dass dieser besagt:

In einem rechtwinkligen Dreieck, in dem die Katheten die Längen  $a$  und  $b$  haben und die Hypotenuse die Länge  $c$ , gilt:

$$a^2 + b^2 = c^2$$

In der Schule geht es hauptsächlich darum, diesen Satz *anzuwenden* (jede\*r kennt wohl irgendwelche mehr oder weniger sinnvollen Textaufgaben mit an die Wände gelehnten Leitern usw.). In der Hochschulmathematik geht es oft jedoch eher um die Frage, *warum* ein bestimmter Satz gültig ist. Mathevorlesungen folgen dabei meist dem Muster „Definition – Satz – Beweis“.

Eine Definition legt erst einmal fest, über welche Objekte man sich unterhält. In einer Definition wird also einem bestimmten Begriff oder Symbol (z.B. den reellen Zahlen  $\mathbb{R}$ ) eine bestimmte *Bedeutung* zugewiesen. Für den Satz des Pythagoras müsste so erst einmal definiert werden

1. Was ist ein Dreieck?
2. Was ist ein rechtwinkliges Dreieck?
3. Was ist überhaupt eine Länge (eine positive reelle Zahl, aber was genau verstehen wir darunter?)

Die Hochschulmathematik ist sehr viel formaler und präziser als die euch bekannte Schulmathematik, z.B. genügt es nicht einfach zu sagen, dass eine reelle Zahl ein gewisser Punkt auf einem „Zahlenstrahl“ (was soll das formal denn sein?) ist. Vielmehr muss ganz genau gesagt werden, welche Eigenschaften so eine reelle Zahl erfüllen muss.

Hat man dann definiert, mit was für Objekten man sich beschäftigen möchte, fragt man sich, was für *Aussagen* aus der Definition der Objekte folgen. Folgt eine Aussage, gibt es also einen Beweis, dass die Aussage immer gilt, so spricht man von einem *Satz*. In unserem Beispiel formulieren wir die Aussage:

## 1. Definition – Satz – Beweis

In einem rechtwinkligen Dreieck, in dem die Katheten die Längen  $a$  und  $b$  haben und die Hypotenuse die Länge  $c$ , gilt:

$$a^2 + b^2 = c^2$$

Dies ist zunächst nur eine Behauptung, die wir aufgestellt haben. Die mathematisch interessante Frage ist: Wieso gilt diese Aussage denn für alle möglichen rechtwinkligen Dreiecke? Die Herleitung der Antwort auf diese Frage wird als *Beweis* bezeichnet. In einem Beweis werden verschiedene Regeln angewendet, um aus bereits bewiesenen Aussagen neue Aussagen herzuleiten. Die Beweisführung wird oft durch das Zeichen  $\square$  oder  $\triangleleft$  bzw. q.e.d.<sup>1</sup> beendet.

## Mathematische Sprache

Es stellt sich natürlich die Frage, wie all dies formuliert wird. Ihr werdet schnell feststellen, dass in der Mathematik häufig eine Mischung aus natürlicher Sprache und mathematischen Symbolen anzutreffen ist, z.B.

Für alle  $a, b \in \mathbb{N}$  gilt: Wenn  $a \leq b$  und  $a \geq b$ , dann ist  $a = b$ .

Diese Art der Notation ist am Anfang vielleicht etwas unverständlich und wenig intuitiv. Allgemein kann man sagen: Lest eine Aussage Stück für Stück, „überfliegt“ sie also nicht und lasst euch, insbesondere am Anfang, nicht abschrecken. In einer einzigen Zeile kann in einem Matheskript schon viel mehr Inhalt stecken als in einem ganzen Absatz voll mit Text.

In der Mathematik haben sich verschiedene Bezeichnungen durchgesetzt, wir versuchen diese hier noch einmal kurz zusammen zu fassen:

**Definition** Legt fest, welche Bedeutung ein Begriff oder Symbol hat.

**Aussage** Eine Behauptung über bestimmte Objekte.

**Beweis** Die Herleitung, dass eine bestimmte Aussage wahr ist.

**Satz** Ein Aussage, für die ein gültiger Beweis existiert.

**Axiom** Eine Aussage, die als wahr angenommen wird und nicht weiter bewiesen werden kann. Beispiel: Für natürliche Zahlen gilt die Annahme, dass jede natürliche Zahl genau eine weitere natürliche Zahl als Nachfolger hat. Axiome stellen somit in gewisser Weise die Basis dar, auf der alle Beweise aufbauen.

**Lemma** Ein Satz, der hauptsächlich dazu verwendet wird, um einen anderen Satz zu beweisen (wird deshalb oft auch als Hilssatz bezeichnet). Die Unterscheidung zwischen Satz und Lemma ist dabei eher subjektiv und liegt im Auge des\*der Autors\*Autorin. Oftmals arbeitet man auf einen „interessanten“ Satz hin und

---

<sup>1</sup>quod erat demonstrandum: lat. für „was zu zeigen war“

beweist auf dem Weg dorthin verschiedene Hilfssätze, um den eigentlichen Beweis möglichst kurz zu halten.

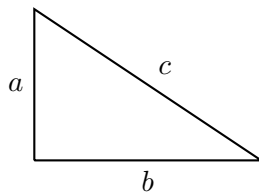
**Korollar** Eine Aussage, die sich aus einem bereits bewiesenen Satz bzw. aus der Definition ohne viel Aufwand beweisen lässt. Dies sind oft Spezialfälle: Weiß man bereits, dass jede durch 4 teilbare Zahl auch gerade ist, weiß man auch direkt, dass jede durch 8 teilbare Zahl gerade ist, da jede durch 8 teilbare Zahl auch durch 4 teilbar ist.

Bei der Durcharbeitung dieses Skripts werdet ihr schon auf einige mathematische Formulierungen und Beweise stoßen; und es gibt ein eigenes Kapitel „Logik und Beweise“.

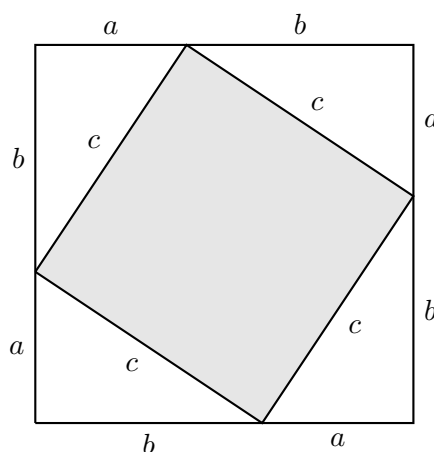
## Beweis: Satz des Pythagoras

Zur Illustration gibt es auch noch einen „Beweis“ des Satz des Pythagoras. Hier sieht man auch gleich, dass man versucht sehr genau in der Beweisführung zu sein, jedoch teilweise auch auf optisches Denkvermögen und „gesunden Menschenverstand“ setzt. Ganz formal gesehen ist das natürlich nicht erlaubt, es kommt jedoch auch immer auf den\*die Dozent\*in und den Kontext an. Für den Beweis setzen wir zudem Grundrechenregeln für die reellen Zahlen sowie den Flächeninhalt eines Quadrats voraus.

*Beweis.* Es seien  $a, b, c \in \mathbb{R}$  wie folgt die Längen in einem rechtwinkligen Dreieck:



Nun betrachten wir ein Quadrat mit der Seitenlänge  $a + b$ , in dieses können wir vier Dreiecke obiger Form wie folgt platzieren:



### 1. Definition – Satz – Beweis

Wir wissen, dass für die Fläche des gesamten Quadrats  $A$  gilt:

$$A = (a + b) \cdot (a + b)$$

Das graue Quadrat in der Mitte hat eine Fläche von  $c^2$  und jedes der vier Dreiecke von  $\frac{1}{2}ab$  (Hälfte des Rechtecks mit den Seitenlängen  $a$  und  $b$ ). Also muss gelten:

$$A = c^2 + 4 \cdot \left(\frac{1}{2}ab\right) = c^2 + 2ab$$

Wir erhalten durch Gleichsetzen:

$$\begin{aligned}(a + b) \cdot (a + b) &= c^2 + 2ab \\ a^2 + 2ab + b^2 &= c^2 + 2ab \\ a^2 + b^2 &= c^2\end{aligned}$$

□

## Beweis: Erweitertes Distributivgesetz

Abschließend sei noch ein Beispiel eines klassischen „Umform-Beweises“ gezeigt: Es ist bekannt, dass das Distributivgesetz gilt:

$$a \cdot (b + c) = ab + ac$$

Wir möchten folgern, dass das erweiterte Distributivgesetz gilt, also dass  $a \cdot (b + c + d) = ab + ac + ad$ . Dazu nehmen wir an, dass das Distributivgesetz bewiesen wurde und damit für unseren neuen Beweis genutzt werden kann. Auch die Assoziativität der Addition setzen wir voraus. Bereits damit lässt sich das Gewünschte beweisen:

*Beweis.*

$$\begin{aligned}a \cdot (b + c + d) &\stackrel{\text{Ass.}}{=} a \cdot (b + (c + d)) \\ &\stackrel{\text{Distr.}}{=} ab + a \cdot (c + d) \\ &\stackrel{\text{Distr.}}{=} ab + ac + ad\end{aligned}$$

□

## 2. Mengenlehre

Aus dem Paradies, das Cantor uns geschaffen hat,  
soll uns niemand vertreiben können.

---

David Hilbert

Der Begriff der Menge ist ein elementarer Bestandteil der Mathematik. Wir folgen der Definition von *Georg Cantor*<sup>1</sup>:

**Definition 2.1** (Menge). Unter einer Menge verstehen wir die Zusammenfassung bestimmter, wohlunterschiedener Objekte unserer Anschauung oder unseres Denkens zu einem Ganzen.<sup>2</sup>

Um eine solche Zusammenfassung zu erhalten, muss also klar sein, welche Objekte zusammengefasst werden. Eine Menge kann dann dadurch beschrieben werden, dass man alle Elemente einfach aufzählt. Die entsprechenden Elemente werden dabei zwischen den Mengenklammern  $\{$  und  $\}$  eingeschlossen und durch Kommata getrennt:

$$A_1 := \{1, 2, 3\}$$

**Hinweis.** Man verwendet  $:=$  um eine *Definition* darzustellen.  $A_1 := \{1, 2, 3\}$  bedeutet also, dass  $A_1$  definiert ist als die folgende Menge.

Man verwendet auch Punkte, um eine Menge zu beschreiben, wenn klar ist, wie die Folge fortzuführen ist:

$$A_2 := \{a, b, c, \dots, z\}$$

Eine andere Möglichkeit, die Elemente einer Menge zu beschreiben, ist über eine *charakteristische Eigenschaft*:

$$A_3 := \{x \mid x \text{ ist eine natürliche Zahl und teilbar durch } 2\}$$

Diese Menge beinhaltet alle geraden natürlichen Zahlen, ihre explizite Notation wäre also  $\{0, 2, 4, 6, \dots\}$ .

Die Menge, die keine Elemente enthält, wird *leere Menge* genannt und durch das Zeichen  $\emptyset$  oder  $\{\}$  symbolisiert. Mengen können selbst wieder Mengen erhalten. So ist z.B.  $\{\emptyset\}$  die Menge, die die leere Menge enthält, und nicht etwa die leere Menge selbst.

---

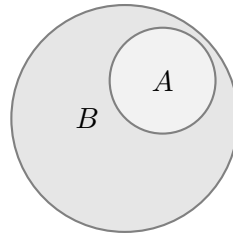
<sup>1</sup>deutscher Mathematiker, 1845–1918

<sup>2</sup>Die hier verwendete Definition einer Menge ist nicht widerspruchsfrei, für viele Anwendungen ist sie jedoch ausreichend.

## 2. Mengenlehre

**Definition 2.2.** Gehört ein Element  $a$  einer Menge  $A$  an, so schreibt man  $a \in A$  (Sprechweise:  $a$  ist Element von  $A$ ). Gehört  $a$  nicht zu  $A$ , so schreibt man  $a \notin A$  (Sprechweise:  $a$  ist nicht Element von  $A$ ).

**Definition 2.3** (Teilmenge). Ist jedes Element einer Menge  $A$  auch in einer Menge  $B$  enthalten, so gilt  $A \subseteq B$  (Sprechweise:  $A$  ist Teilmenge von  $B$ ). Man sagt auch  $B \supseteq A$  (Sprechweise:  $B$  ist Obermenge von  $A$ ).



**Abbildung 2.1.:** Ist  $A \subseteq B$ , so liegt jedes Element von  $A$  auch in  $B$ .  $B$  kann jedoch auch mehr Elemente als  $A$  enthalten.

**Definition 2.4** (Gleichheit von Mengen). Zwei Mengen  $A$  und  $B$  sind genau dann gleich, wenn sie beide die gleichen Elemente enthalten, wenn also gilt  $A \subseteq B$  und  $B \subseteq A$ .

Damit ist klar, dass eine Menge ein Element nicht mehrfach enthalten kann. So gilt zum Beispiel  $\{a, a, b\} = \{a, b\}$ .

**Definition 2.5** (Echte Teilmenge). Ist jedes Element einer Menge  $A$  auch in einer Menge  $B$  enthalten und es gibt mindestens ein Element in  $B$ , welches nicht in  $A$  ist, so gilt  $A \subset B$  (Sprechweise:  $A$  ist echte Teilmenge von  $B$ ). Man sagt auch  $B \supset A$  (Sprechweise:  $B$  ist echte Obermenge von  $A$ ).

**Hinweis.** Man schreibt auch  $A \subsetneq B$  bzw.  $B \supsetneq A$  für  $A \subset B$  bzw.  $B \supset A$ .

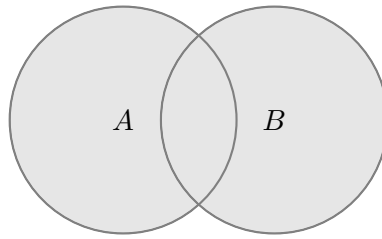
### 2.1. Mengenoperationen

Mengen können durch verschiedene Operationen zu neuen Mengen verknüpft werden. Diese Operationen lassen sich mit *Venn-Diagrammen* darstellen (benannt nach *John Venn*<sup>3</sup>). Im Folgenden seien  $A$  und  $B$  zwei beliebige Mengen.

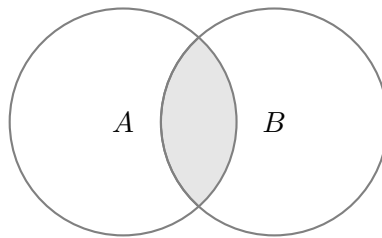
- Die *Vereinigung*, geschrieben  $A \cup B$ , gesprochen  $A$  vereinigt  $B$ , enthält alle Elemente, die in  $A$  oder in  $B$  enthalten sind:  $A \cup B := \{x \mid x \in A \text{ oder } x \in B\}$ .

---

<sup>3</sup>englischer Mathematiker, 1834 - 1923

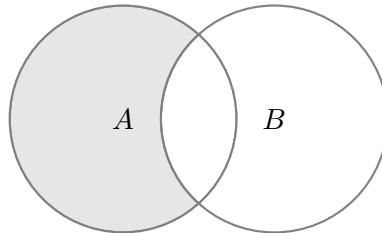


- Der *Durchschnitt*, geschrieben  $A \cap B$ , gesprochen *A geschnitten B*, enthält alle Elemente, die in *A* und in *B* enthalten sind:  $A \cap B := \{x \mid x \in A \text{ und } x \in B\}$ .

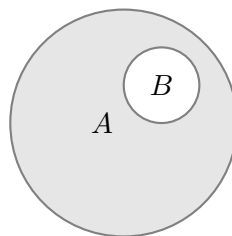


Ist  $A \cap B = \emptyset$ , so nennt man *A* und *B* *disjunkt*.

- Die *Differenz*, geschrieben  $A \setminus B$ , gesprochen *A ohne B*, enthält alle Elemente aus *A*, die *nicht* in *B* enthalten sind:  $A \setminus B := \{x \mid x \in A \text{ und } x \notin B\}$ .



Ist  $B \subseteq A$ , so nennt man  $A \setminus B$  auch das *Komplement von B in A* und schreibt  $B^C$ , gesprochen *B Komplement*, (wenn man weiß, auf welche Menge sich das Komplement bezieht).



## 2. Mengenlehre

**Beispiel 2.6** (Mengenoperationen). Sei

- $A = \{\blacktriangle, \blacktriangledown, \star\}$
- $B = \{\dagger, \clubsuit\}$
- $C = \{\blacktriangle, \star, \clubsuit\}$
- $\mathbb{N}_g = \{n \in \mathbb{N} \mid n \text{ ist gerade}\} = \{0, 2, 4, \dots\}$

$$A \cup B = \{\blacktriangle, \blacktriangledown, \star, \dagger, \clubsuit\}$$

$$A \cup C = \{\blacktriangle, \blacktriangledown, \star, \clubsuit\}$$

$$A \cap B = \emptyset$$

$$A \cap C = \{\blacktriangle, \star\}$$

$$A \setminus B = \{\blacktriangle, \blacktriangledown, \star\} = A$$

$$A \setminus C = \{\blacktriangledown\}$$

$$\mathbb{N} \setminus \mathbb{N}_g = \mathbb{N}_g^C = \{n \in \mathbb{N} \mid n \text{ ist ungerade}\} = \{1, 3, 5, \dots\}$$

**Satz 2.7** (Eigenschaften der Mengenoperationen). Für die Mengenoperationen gelten die folgenden Eigenschaften:

**Assoziativität**  $A \cup (B \cup C) = (A \cup B) \cup C$  bzw.  $A \cap (B \cap C) = (A \cap B) \cap C$

**Distributivität**  $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$  bzw.  
 $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$

**Kommutativität**  $A \cup B = B \cup A$  bzw.  $A \cap B = B \cap A$

**Idempotenz**  $A \cup A = A$  bzw.  $A \cap A = A$

*Beweis.* Wir beweisen exemplarisch die Assoziativität.

Um zu zeigen, dass  $A \cup (B \cup C) = (A \cup B) \cup C$ , müssen wir zeigen, dass  $A \cup (B \cup C) \subseteq (A \cup B) \cup C$  und  $A \cup (B \cup C) \supseteq (A \cup B) \cup C$ .

1.  $\subseteq$  Es sei  $x \in A \cup (B \cup C)$ . Also  $x \in A$  oder  $x \in B \cup C$ , also  $x \in A$  oder  $x \in B$  oder  $x \in C$ . Dann ist  $x \in A \cup B$  oder  $x \in C$ . Also  $x \in (A \cup B) \cup C$ .
2.  $\supseteq$  Es sei  $x \in (A \cup B) \cup C$ , also  $x \in A \cup B$  oder  $x \in C$ , also  $x \in A$  oder  $x \in B$  oder  $x \in C$ . Dann ist  $x \in A$  oder  $x \in B \cup C$ . Also  $x \in A \cup (B \cup C)$ .

□

Die Regeln lassen sich einfacher beweisen, wenn man mit den Regeln der *Aussagenlogik* (Kapitel 6) argumentiert.

**Definition 2.8** (Kardinalität). Die Kardinalität, geschrieben  $|A|$ , einer Menge  $A$  gibt an, wie viele Elemente in  $A$  enthalten sind. Enthält  $A$  unendlich viele Elemente, so schreibt man  $|A| = \infty$ .

**Beispiel 2.9** (Kardinalität von Mengen).

- $|\emptyset| = 0$
- $|\{x \mid x \text{ ist eine gerade Primzahl}\}| = |\{2\}| = 1$
- $|\{a, b\}| = 2$
- $|\{\heartsuit, \spadesuit, \diamondsuit\}| = 3$
- $|\mathbb{N}| = \infty$

**Definition 2.10** (Potenzmenge). Die *Potenzmenge* einer Menge  $A$  ist die Menge aller Teilmengen von  $A$  und wird mit  $\mathcal{P}(A)$  bezeichnet. Formal:

$$\mathcal{P}(A) := \{T \mid T \subseteq A\}$$

**Hinweis.** Manchmal wird die Potenzmenge von  $A$  auch geschrieben als  $2^A$ .

**Beispiel 2.11.** Die Potenzmenge der Menge  $A = \{1, 2, 3\}$  ist

$$\mathcal{P}(A) = \{\emptyset, \{1\}, \{2\}, \{3\}, \{1, 2\}, \{1, 3\}, \{2, 3\}, \{1, 2, 3\}\}$$

**Definition 2.12** (Kartesisches Produkt). Als *Kartesisches Produkt* oder *Produktmenge*  $A \times B$  (gesprochen  $A$  Kreuz  $B$ ) von zwei Mengen bezeichnet man die Menge aller Paare  $(a, b)$ , wobei  $a$  aus  $A$  und  $b$  aus  $B$  ist. Formal:

$$A \times B := \{(a, b) \mid a \in A \text{ und } b \in B\}$$

**Hinweis.** Bei Paaren ist die Reihenfolge der Elemente *nicht egal*. Bei Mengen spielt es keine Rolle, in welcher Reihenfolge man die Elemente aufzählt. So ist zum Beispiel  $\{1, 2\} = \{2, 1\}$ . Bei Paaren gilt jedoch  $(1, 2) \neq (2, 1)$ .

**Beispiel 2.13.** Das Kartesische Produkt der Mengen  $A = \{J, Q, K\}$  und  $B = \{\heartsuit, \spadesuit\}$  ist

$$A \times B = \{(J, \heartsuit), (J, \spadesuit), (Q, \heartsuit), (Q, \spadesuit), (K, \heartsuit), (K, \spadesuit)\}$$

## 2.2. Relationen

Mathematicians do not study objects, but relations between objects. Thus, they are free to replace some objects by others so long as the relations remain unchanged.

---

*Henri Poincaré*

Wie wir gelernt haben, sind Mengen eine Zusammenfassung von verschiedenen Objekten. Außerdem haben wir einige Operationen auf Mengen kennengelernt. Oft benötigt man aber noch Beziehungen zwischen Objekten, die wir *Relationen* nennen. Relationen, die bereits aus der Schule bekannt sein sollten, sind beispielsweise die Gleichheitsrelation  $=$  und die Ungleichheitsrelation  $\neq$ . Diese beiden sind *binäre Relationen*, da sie angeben, ob *zwei* Elemente zusammen die Eigenschaft „Gleichheit“ bzw. „Ungleichheit“ erfüllen.

## 2. Mengenlehre

**Definition 2.14** (Binäre Relation). Seien  $A, B$  Mengen. Eine Teilmenge  $R \subseteq A \times B$  nennt man *binäre Relation* zwischen  $A$  und  $B$ . Anstatt  $(a, b) \in R$  schreibt man oft auch  $aRb$  (*Infix-Schreibweise*).

Für einige häufig benutzte Relationen gibt es eigene Symbole, so zum Beispiel die aus der Schule bekannten Relationen  $<, \leq, >, \geq$ .

**Beispiel 2.15** (Gleichheitsrelation). Die Gleichheitsrelation  $=$  über der Menge der natürlichen Zahlen ist eine Teilmenge von  $\mathbb{N} \times \mathbb{N}$  und gegeben durch  $\{(0, 0), (1, 1), (2, 2), (3, 3), \dots\}$ . Man schreibt  $(42, 42) \in =$  oder  $42 = 42$ .

**Beispiel 2.16** (Kleinerrelation). Die Kleinerrelation  $<$  über der Menge der natürlichen Zahlen ist eine Teilmenge von  $\mathbb{N} \times \mathbb{N}$  und gegeben durch  $\{(0, 1), (0, 2), (1, 2), (0, 3), (1, 3), \dots\}$ . Formal lässt sich die Relation wie folgt beschreiben:

$$< := \{(a, b) \mid a, b \in \mathbb{N} \text{ und es gibt ein } c \in \mathbb{N} \text{ mit } c \neq 0, \text{ sodass } a + c = b\}$$

Man schreibt  $(21, 42) \in <$  oder  $21 < 42$ .

Es existieren darüber hinaus auch *unäre*, *ternäre* und allgemein *n-äre* Relationen. Sie sind lediglich Verallgemeinerungen der obigen Definition und werden hier nicht weiter behandelt.

**Beispiel 2.17** (Relation). Sei

$$P := \{\text{Alice, Bob, Carol, Dave}\}$$

$$B := \{\text{Der Proceß, Die Physiker, Andorra}\}$$

$$A := \{\text{Max Frisch, Franz Kafka, Friedrich Dürrenmatt}\}$$

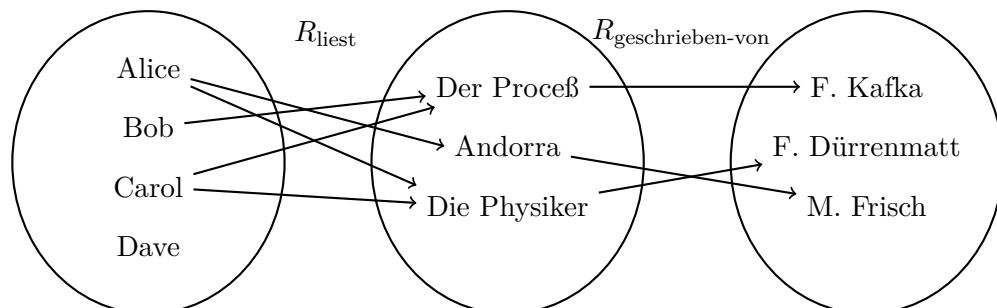
Wir definieren die Relation  $R_{\text{liest}} \subseteq P \times B$  wie folgt:

$$R_{\text{liest}} := \{(\text{Alice, Andorra}), (\text{Alice, Die Physiker}), (\text{Bob, Der Proceß}), (\text{Carol, Der Proceß}), (\text{Carol, Die Physiker})\}$$

und die Relation  $R_{\text{geschrieben-von}} \subseteq B \times A$  wie folgt:

$$R_{\text{geschrieben-von}} := \{(\text{Der Proceß, Franz Kafka}), (\text{Andorra, Max Frisch}), (\text{Die Physiker, Friedrich Dürrenmatt})\}$$

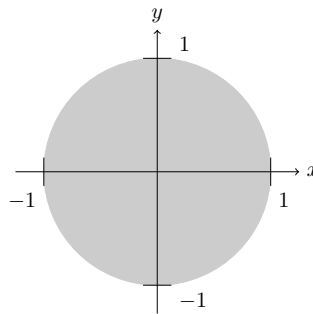
Graphisch lässt sich dies z. B. wie folgt veranschaulichen:



**Beispiel 2.18** (Kreis-Relation). Sei

$$K := \{(x, y) \in \mathbb{R} \times \mathbb{R} \mid x^2 + y^2 \leq 1\}$$

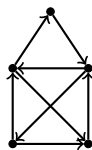
$K$  definiert alle Punkte in einem Kreis um den *Ursprung* mit Radius 1. Siehe Abbildung 2.2.



**Abbildung 2.2.:** Die  $K$ -Relation definiert einen Einheitskreis in der Ebene.

**Beispiel 2.19** (Endlicher, gerichteter Graph). Ein endlicher, gerichteter Graph ist ein Tupel  $G = (V, E)$  wobei  $V$  eine endliche Menge von *Knoten* (engl. *vertices*) und  $E$  eine Relation  $E \subseteq V \times V$  ist. Die Elemente von  $E$  bezeichnet man auch als *Kanten* (engl. *edges*).

Die *Graphentheorie* ist ein Teilgebiet der Mathematik und insbesondere für die Informatik bedeutend, da viele Algorithmen auf Graphen basieren bzw. sich auf Graphen zurückführen lassen. Beispiele für Graphen sind Straßennetze in denen die Knoten für Wegpunkte stehen und die Kanten Wegabschnitte sind, welche Punkte miteinander verbinden. Auch ein soziales Netzwerk kann man sich als einen großen Graph vorstellen: Die Mitglieder bilden die Knoten und die Kanten sind gegeben durch eine friend-of-Relation.



**Abbildung 2.3.:** Ein endlicher, gerichteter Graph besteht aus einer Menge von *Knoten* und *Kanten*, die Knoten miteinander verbinden.

### 2.2.1. Komposition

Wie schon gesagt, repräsentieren Relationen Eigenschaften. Durch Komposition können wir nun die Eigenschaften von zwei bekannten Relationen zu einer neuen kombinieren.

## 2. Mengenlehre

**Definition 2.20** (Komposition von Relationen). Sei  $R \subseteq A \times B$  und  $S \subseteq B \times C$ . Dann ist  $R \circ S \subseteq A \times C$  (gesprochen  $R$  verknüpft mit  $S$  oder  $R$  nach  $S$ ) eine Relation, die wie folgt definiert ist:

$$R \circ S := \{(a, c) \in A \times C \mid \text{es gibt ein } b \in B \text{ mit } (a, b) \in R \text{ und } (b, c) \in S\}$$

**Hinweis.** Manche Autor\*innen verwenden auch die Schreibweise  $S \circ R$  anstatt  $R \circ S$ . Aus den involvierten Mengen sollte sich die Reihenfolge klar ergeben.

**Hinweis.** „Es gibt ein“ bedeutet in der Mathematik, dass es *mindestens* ein Element gibt, für welches die Aussage gilt.

**Beispiel 2.21** (Komposition). Für die Relationen  $R_{\text{liest}}$  und  $R_{\text{geschrieben-von}}$  aus Beispiel 2.17 auf Seite 14 ist  $R_{\text{liest}} \circ R_{\text{geschrieben-von}}$  gegeben durch

$$R_{\text{liest}} \circ R_{\text{geschrieben-von}} = \{(Alice, Max Frisch), (Alice, Friedrich Dürrenmatt), \\ (Bob, Franz Kafka), (Carol, Franz Kafka), \\ (Carol, Friedrich Dürrenmatt)\}$$

Wir können dies auffassen als eine Relation  $R_{\text{liest-Buch-von}}$ .

Es ist z. B.  $(Alice, Max Frisch) \in R_{\text{liest}} \circ R_{\text{geschrieben-von}}$ , da  $(Alice, Andorra) \in R_{\text{liest}}$  und  $(Andorra, Max Frisch) \in R_{\text{geschrieben-von}}$ .

$(Bob, Friedrich Dürrenmatt)$  ist jedoch nicht in  $R_{\text{liest}} \circ R_{\text{geschrieben-von}}$ , da es kein Buch gibt, welches von Bob gelesen wird und von Friedrich Dürrenmatt geschrieben wurde.

### 2.2.2. Eigenschaften von Relationen

Oftmals sind Relationen der Form  $R \subseteq A \times A$  (Sprechweise: Relationen über  $A$ ) von besonderem Interesse. Diese Relationen können bestimmte Eigenschaften haben, von denen wir einige exemplarisch aufführen wollen:

**Definition 2.22** (Eigenschaften von binären Relationen). Sei  $R \subseteq A \times A$  eine Relation über  $A$ . Dann heißt  $R$ :

**reflexiv** falls für alle  $a \in A$  gilt:  $(a, a) \in R$ .

**irreflexiv** falls für alle  $a \in A$  gilt:  $(a, a) \notin R$

**symmetrisch** falls für alle  $a, b \in A$  gilt: Wenn  $(a, b) \in R$ , dann  $(b, a) \in R$ .

**antisymmetrisch** falls für alle  $a, b \in A$  gilt: Wenn  $(a, b) \in R$  und  $(b, a) \in R$ , dann  $a = b$ .

**transitiv** falls für alle  $a, b, c \in A$  gilt: Wenn  $(a, b) \in R$  und  $(b, c) \in R$ , dann  $(a, c) \in R$ .

**Beispiel 2.23.**

- < Die Kleinerrelation ist *irreflexiv* ( $5 \not< 5$ ), *antisymmetrisch* ( $x < y$  und  $y < x$  kommt nicht vor, deswegen gilt die Eigenschaft) und *transitiv* (wenn  $5 < 6$  und  $6 < 7$ , dann  $5 < 7$ ).
- $\geq$  Die Größergleichrelation ist *reflexiv* ( $5 \geq 5$ ), *antisymmetrisch* (wenn  $x \geq y$  und  $y \geq x$ , dann  $x = y$ ) und *transitiv* (wenn  $7 \geq 6$  und  $6 \geq 5$ , dann  $7 \geq 5$ ).
- $K$  aus Beispiel 2.18. Die Relation ist *symmetrisch* (für einen Punkt  $(x, y) \in K$  liegt auch immer der gespiegelte Punkt  $(y, x)$  im Kreis). Die Relation ist nicht *reflexiv*, denn es ist z. B.  $(5, 5) \notin K$ , und auch nicht *transitiv*, denn es ist  $(1, 0) \in K$  und  $(0, 1) \in K$ , aber nicht  $(1, 1)$ . Die Relation ist auch nicht *irreflexiv*, denn  $(0, 0) \in K$ .



## 3. Zahlenmengen und Operationen

Numbers exist only in our minds. There is no physical entity that is number 1. If there were, 1 would be in a place of honor in some great museum of science, and past it would file a steady stream of mathematicians gazing at 1 in wonder and awe.

---

*John B. Fraleigh / Raymond A. Beauregard*

### 3.1. Zahlenräume

#### 3.1.1. Natürliche Zahlen

Die grundlegendste Zahlenmenge ist die Menge der natürlichen Zahlen, die man mit  $\mathbb{N}$  bezeichnet. Natürliche Zahlen können zum (Ab-)Zählen von „Dingen“ verwendet werden: Es sind 5 Kühe auf der Weide oder 42 Personen im Bus. Dabei spielt die Zahl 0 (leider) an dieser Fakultät eine Sonderrolle. Während in vielen Bereichen der Mathematik die Menge  $\mathbb{N}$  die Zahl 0 nicht enthält (und die Menge mit 0 als  $\mathbb{N}_0$  bezeichnet wird), ist es in der Informatik üblich, dass die natürlichen Zahlen stets auch die 0 enthalten. In diesem Skript verwenden wir die natürlichen Zahlen immer mit 0 und definieren also:

**Definition 3.1** (Natürliche Zahlen).  $\mathbb{N} := \{0, 1, 2, 3, \dots\}$

Diese „Definition“ ist für die Mathematik jedoch nicht formal genug. Wir werden also versuchen, die natürlichen Zahlen etwas mathematischer zu fassen.

#### Peano-Axiome

Um grundlegende Regeln der Mathematik zu beschreiben, werden Axiome verwendet. Diese werden nicht bewiesen, sondern wir nehmen sie als gegebenes Fundament für alles Weitere an. Wir wollen also angeben, welche Eigenschaften die natürlichen Zahlen erfüllen müssen.

Bei den natürlichen Zahlen heißen diese Peano-Axiome (nach *Guisepppe Peano*<sup>1</sup>).

**Definition 3.2** (Peano-Axiome).

1. Es gibt ein Element 0, welches eine natürliche Zahl ist, also  $0 \in \mathbb{N}$ .

---

<sup>1</sup>italienischer Mathematiker, 1858–1932

### 3. Zahlenmengen und Operationen

2. Ist  $n$  eine natürliche Zahl, so hat sie eine natürliche Zahl  $n^+$  als Nachfolger, also wenn  $n \in \mathbb{N}$ , dann gilt  $n^+ \in \mathbb{N}$ .
3. Keine natürliche Zahl hat 0 als Nachfolger, also wenn  $n \in \mathbb{N}$ , dann gilt  $n^+ \neq 0$ .
4. Haben zwei natürliche Zahlen den gleichen Nachfolger, so sind sie gleich, also wenn  $n, m \in \mathbb{N}$  und  $n^+ = m^+$ , dann gilt  $n = m$ .

Hiermit soll festgelegt werden, dass jede Zahl genau einen eindeutigen Nachfolger hat.

5. Sind für eine beliebige Menge  $X$  die ersten drei Axiome erfüllt (also enthält sie 0 und für jede natürliche Zahl deren Nachfolger), so ist  $X$  eine Teilmenge von  $\mathbb{N}$ .

Dieses Axiom legt die Eindeutigkeit der Menge der natürlichen Zahlen fest. Wenn die Menge  $X$  also die Axiome erfüllt, darf man davon ausgehen, dass für  $X$  alle Eigenschaften der natürlichen Zahlen gelten.

Im Folgenden wollen wir noch einmal auf wichtige Rechenoperationen eingehen.

**Hinweis.** Wir gehen einmal davon aus, dass jede\*r Leser\*in weiß, wie man natürliche Zahlen miteinander addiert. Vielmehr ist es Ziel des folgenden Abschnitts zu zeigen, dass all diese Konzepte einen formalen Hintergrund haben und wie man einfache Dinge wie Addition sauber in einem mathematischen Modell definieren kann. Zudem soll der Abschnitt dazu dienen, grundlegende Rechenregeln zu wiederholen.

**Definition 3.3** (Addition).

$$\begin{aligned}n + 0 &= n \\n + m^+ &= (n + m)^+\end{aligned}$$

Eine intuitive Interpretation der zweiten Zeile lautet, dass man für die Addition zweier natürlichen Zahlen  $n$  und  $m$  einfach den  $m$ -ten Nachfolger von  $n$  nimmt. So ist zum Beispiel  $5 + 2 = (5 + 1)^+ = ((5 + 0)^+)^+$ . Bei dieser Definition handelt es sich um eine sogenannte *rekursive Definition*, diese werden wir in Abschnitt 4.2.4 noch genauer behandeln.

Die Addition ist *kommutativ* (es gilt also  $2 + 3 = 3 + 2$ ) und *assoziativ* (es gilt also  $2 + (3 + 4) = (2 + 3) + 4$ )

Genauso wie die Interpretation der Addition als wiederholtes Anwenden des Nachfolge-Operators, können wir die Multiplikation als wiederholte Addition verstehen:

$$a \cdot b := \underbrace{a + a + a + \cdots + a}_{b\text{-mal}}$$

Auch die Multiplikation ist *kommutativ* und *assoziativ*. Addition und Multiplikation sind zusammen *distributiv*, also  $a \cdot (b + c) = a \cdot b + a \cdot c$ .

Die Potenzierung können wir wiederum als wiederholtes Multiplizieren darstellen:

$$a^b := \underbrace{a \cdot a \cdot a \cdot \dots \cdot a}_{b\text{-mal}}$$

Im Gegensatz zur Addition und Multiplikation ist die Potenzierung nicht *kommutativ* ( $2^3 = 8 \neq 9 = 3^2$ ) und auch nicht *assoziativ* ( $(3^1)^3 = 27 \neq 3 = 3^{(1^3)}$ ).

### 3.1.2. Ganze Zahlen

**Definition 3.4** (Subtraktion). Parallel zur Addition existiert die Subtraktion.

$$a - b = c \text{ gdw. }^2 a = b + c$$

Wie man leicht sehen kann, ist diese Operation auf den natürlichen Zahlen nicht abgeschlossen, so gibt es zum Beispiel keine natürliche Zahl  $c$ , die die Gleichung  $2 - 5 = c$  bzw.  $5 + c = 2$  erfüllt. Um diese Operation abzuschließen, muss der Zahlenraum  $\mathbb{N}$  also auf einen größeren Zahlenraum  $\mathbb{Z} \supset \mathbb{N}$  erweitert werden. Das wird durch Hinzufügen des Vorzeichens  $-$  erreicht. Auf die Erklärung, wie das Vorzeichen in den bisher vorgestellten Operationen behandelt wird, werden wir nicht eingehen, dies sollte bestens bekannt sein. Die Menge der ganzen Zahlen ist also

$$\mathbb{Z} := \{0, 1, -1, 2, -2, \dots\}$$

### 3.1.3. Rationale Zahlen

**Definition 3.5** (Division). Parallel zur Multiplikation existiert die Division.

$$\frac{a}{b} = c \text{ gdw. } a = b \cdot c$$

Auch hier ist leicht ein Beispiel zu finden, bei dem unser bisheriger Zahlenraum  $\mathbb{Z}$  nicht ausreicht, wie  $\frac{5}{2} = c$  bzw.  $2 \cdot c = 5$ . Um auch Operationen dieser Art durchführen zu können, erweitern wir den Zahlenraum auf die *rationalen Zahlen*  $\mathbb{Q} \supset \mathbb{Z}$ .

**Definition 3.6** (Rationale Zahlen).

$$\mathbb{Q} := \left\{ \frac{a}{b} \mid a, b \in \mathbb{Z}, b \neq 0 \right\}$$

Bei dieser Definition ist die Zahlendarstellung, im Gegensatz zu  $\mathbb{Z}$ , nicht eindeutig. Als Beispiel gäbe es hier:  $\frac{6}{24} = \frac{2}{8} = \frac{3}{12} = \frac{-6}{-24} = \dots$  Zusätzlich kann man die rationalen Zahlen auch in Dezimaldarstellung schreiben, wobei sich unendlich lange, jedoch periodische Nachkommastellen ergeben können:  $\frac{4}{7} = 0.5714285\overline{7}$

**Hinweis.** Wir verwenden in diesem Skript die englische Zahlennotation, also den Punkt als Dezimaltrennzeichen und das Komma als Tausendertrennzeichen.

<sup>2</sup>"Genau dann, wenn"; die erste Aussage gilt, wenn die zweite Aussage gilt, und nur dann. Beide Aussagen sind also Äquivalent.

### 3. Zahlenmengen und Operationen

#### 3.1.4. Reelle Zahlen

Der Grund, eine weitere Zahlenmenge einzuführen, ist, dass es auch Zahlen gibt, die eine unendlich lange, nicht-periodische Dezimaldarstellung haben und sich damit nicht als rationale Zahlen darstellen lassen, sondern sogenannte *irrationale* Zahlen sind. Bekannte Beispiele sind die Zahlen  $\sqrt{2}$  und  $\pi$ . Die Menge der reellen Zahlen wird mit  $\mathbb{R}$  bezeichnet. Auch hier wollen wir auf eine formale Definition verzichten.

**Definition 3.7** (Radizieren). Die Lösung einer Gleichung der Form  $a^n = x$  ist die  $n$ -te *Wurzel* von  $x$ . Man schreibt also

$$\sqrt[n]{x} = a \text{ gdw. } a^n = x$$

Hierbei nennt man das Ergebnis von  $\sqrt[n]{x}$  *Wurzel* oder *Radix*,  $n$  ist der *Wurzelexponent* und  $x$  bezeichnet man als *Radikand* oder *Wurzelbasis*.

**Hinweis.** Aus dieser Definition folgt auch, dass die *Lösungsmenge* von  $\sqrt[2]{4} = c$  die Zahlen 2 und  $-2$  enthält.

**Definition 3.8** (Logarithmieren). Die Lösung einer Gleichung der Form  $b^x = a$  ist der *Logarithmus zur Basis*  $b$  von  $a$ . Man schreibt

$$\log_b a = x \text{ gdw. } b^x = a$$

Hier heißt  $a$  *Numerus*,  $b$  ist die *Basis*.

**Definition 3.9** (Betrag). Die Betragsfunktion  $|x|$  ordnet der Zahl  $x$  den Abstand zur Null zu. Sie ist stets eine positive reelle Zahl (oder Null) und definiert durch

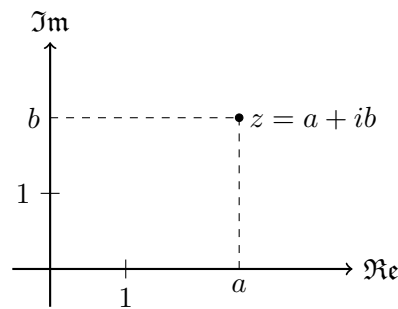
$$|x| = \begin{cases} x & \text{für } x \geq 0 \\ -x & \text{für } x < 0 \end{cases}$$

Insbesondere ist damit für  $a, b \in \mathbb{R}$  der Abstand der beiden Zahlen zueinander  $|a - b| = |b - a|$ .

**Hinweis.** Wir verwenden hier eine sogenannte *Fallunterscheidung*. Man verwendet hierfür eine geschweifte Klammer. Auf der rechten Seite stehen die verschiedenen möglichen Fälle und auf der linken Seite die entsprechenden Werte.

#### 3.1.5. Komplexe Zahlen

Die reellen Zahlen sind nicht abgeschlossen bezüglich dem Radizieren. Um beispielsweise eine Lösung für  $c = \sqrt{-1}$  zu finden, müsste für  $c$  gelten, dass  $c^2 = -1$ . Eine solche Zahl existiert jedoch in  $\mathbb{R}$  nicht, denn für jede Zahl  $c \in \mathbb{R}$  ist offensichtlich  $c^2 \geq 0$ , also  $\sqrt{-1} \notin \mathbb{R}$ .



**Abbildung 3.1.:** Eine komplexe Zahl  $z = a + ib$  lässt sich in einem zweidimensionalen Koordinatensystem darstellen. Auf der  $\Re$ -Achse wird der reelle Anteil der Zahl abgetragen und auf der  $\Im$ -Achse der imaginäre Teil.

Als Lösung wird der Zahlenraum noch einmal zu den sogenannten *komplexen Zahlen* erweitert, welche mit  $\mathbb{C}$  bezeichnet werden. Wir fordern, dass jede reelle Zahl auch eine komplexe Zahl ist, dass also gilt  $\mathbb{R} \subset \mathbb{C}$ .

Die Idee ist, den Zahlenstrahl wie bei den bisherigen Erweiterungen der natürlichen Zahlen, so zu erweitern, dass er mehr Zahlen darstellen kann. Der reelle Zahlenstrahl (eindimensional) lässt sich allerdings nicht mehr entlang seiner Achse erweitern, weshalb nur noch eine Erweiterung in die zweite Dimension bleibt (siehe Abbildung 3.1). Zu diesem Zweck definieren wir einfach  $i^2 := -1$ , also  $i = \sqrt{-1}$  und bezeichnen  $i$  als die *imaginäre Einheit*. Wie bereits gesagt ist  $i$  keine reelle Zahl.

Eine komplexe Zahl besteht nun aus zwei Komponenten: Einem reellen Teil und einem komplexen Teil.

**Definition 3.10** (Komplexe Zahlen). Die Menge der komplexen Zahlen,  $\mathbb{C}$ , ist wie folgt definiert:

$$\mathbb{C} := \{a + i \cdot b \mid a, b \in \mathbb{R}\}$$

Eine reelle Zahl  $x \in \mathbb{R}$  können wir offensichtlich als komplexe Zahl darstellen, wobei der imaginäre Teil der Zahl 0 ist.

### Rechenregeln

Aus der Definition von  $i$  folgt zunächst, dass

$$i \cdot i = \sqrt{-1} \cdot \sqrt{-1} = -1$$

Für Addition und Multiplikation von komplexen Zahlen kann man bei dieser Schreibweise einfach ignorieren, dass es sich um komplexe Zahlen handelt. Betrachtet man  $i$  einfach als beliebige Konstante, kann man damit rechnen wie aus der Schule bekannt und  $i$  jeweils ausmultiplizieren beziehungsweise ausklammern. Für die Addition zweier komplexer Zahlen  $a + bi$  und  $c + di$  gilt demnach

$$(a + bi) + (c + di) = (a + c) + i(b + d)$$

### 3. Zahlenmengen und Operationen

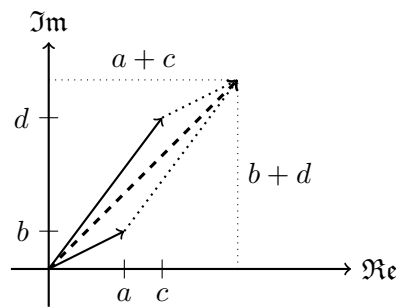


Abbildung 3.2.: Addition von  $(a + i \cdot b) + (c + i \cdot d)$ .

Die Addition zweier komplexer Zahlen können wir als das Aneinandersetzen der Pfeile, die vom Ursprung zum jeweiligen Punkt auf der komplexen Ebene reichen, auffassen (siehe Abbildung 3.2).

Ganz analog funktioniert auch die Subtraktion.

Für die Multiplikation gilt entsprechend

$$(a + bi)(c + di) = (ac - bd) + i(ad + bc)$$

Für die Division gilt

$$\frac{(a + bi)}{(c + di)} = \frac{(a + bi)(c - di)}{(c + di)(c - di)} = \frac{(a + bi)(c - di)}{(c^2 + d^2)}$$

#### Polarform

Eine komplexe Zahl haben wir interpretiert als einen Punkt in der zweidimensionalen Ebene. Eine weitere Möglichkeit eine komplexe Zahl zu charakterisieren ist die sogenannte *Polarform*. In dieser Form wird die Länge des Zeigers  $r$  und ein Winkel  $0 \leq \varphi \leq 2\pi$  angegeben (mit der formalen Definition von Winkeln sowie den Funktionen *sin*, *cos* beschäftigen wir uns noch in Kapitel 4.4.4).

Eine komplexe Zahl  $z$  in Polarform ist dann gegeben durch

$$z = r(\cos \varphi + i \cdot \sin \varphi)$$

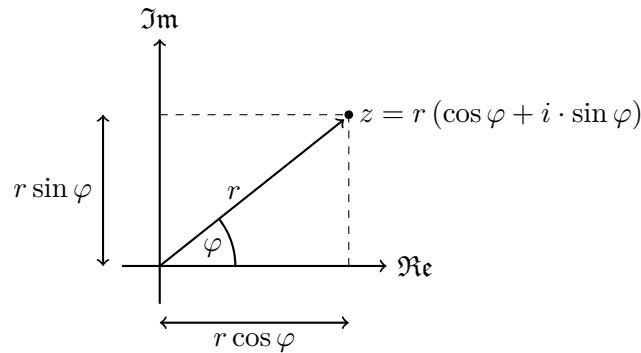
Graphisch ist das in Abbildung 3.3 dargestellt.

Unter Verwendung der *Eulerschen Formel*<sup>3</sup> lassen sich komplexe Zahlen in Polarform auch anders darstellen:

$$z = r(\cos \varphi + i \cdot \sin \varphi) = r \cdot e^{i\varphi}$$

Diese Darstellung ist unter anderem in der Elektrotechnik gebräuchlich, wenn Wechselspannungssignale beschrieben werden.

<sup>3</sup>Leonard Euler, schweizer Mathematiker, 1707 - 1783



**Abbildung 3.3.:** Polarform einer komplexen Zahl in der komplexen Zahlenebene.

**Satz 3.11** (Umwandlung verschiedener Darstellungsformen). Gegeben sei eine komplexe Zahl  $z = a + i \cdot b$ . Dann ist die Polarform von  $z$  gegeben durch:

$$r = \sqrt{a^2 + b^2}$$

$$\varphi = \begin{cases} \arctan \frac{b}{a} & \text{für } a > 0 \\ \arctan \frac{b}{a} + \pi & \text{für } a < 0, b \geq 0 \\ \arctan \frac{b}{a} - \pi & \text{für } a < 0, b < 0 \\ \frac{\pi}{2} & \text{für } a = 0, b > 0 \\ -\frac{\pi}{2} & \text{für } a = 0, b < 0 \\ 0 & \text{für } a = 0, b = 0 \end{cases}$$

Gegeben sei die komplexe Zahl  $z$  in Polarform mit Zeigerlänge  $r$  und Winkel  $\varphi$ . Dann ist  $z = a + i \cdot b$  mit

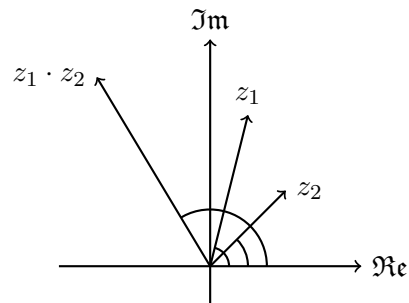
$$a = \Re(z) = r \cdot \cos \varphi$$

$$b = \Im(z) = r \cdot \sin \varphi$$

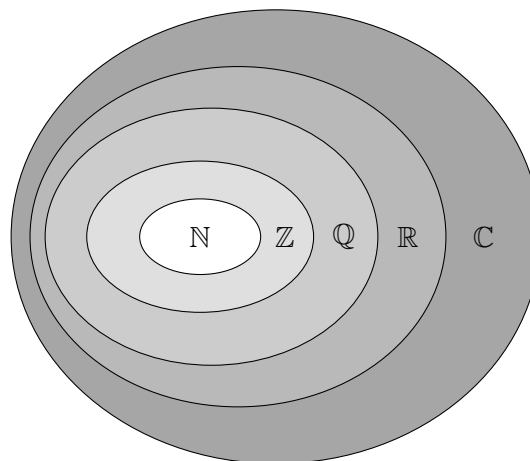
**Hinweis.** Für zwei komplexe Zahlen in der Polarkoordinatendarstellung bedeutet die Multiplikation, dass sich die Winkel der Zahlen *addieren* und die Längen multiplizieren, siehe Abbildung 3.4.

Abschließend werden in Abbildung 3.5 die Teilmengenbeziehungen zwischen den behandelten Zahlenmengen dargestellt.

### 3. Zahlenmengen und Operationen



**Abbildung 3.4.:** Multiplikation zweier komplexer Zahlen in Polarkoordinatendarstellung.



**Abbildung 3.5.:** Teilmengenbeziehungen zwischen den Zahlenmengen (hier in Zwiebeldarstellung):  $\mathbb{N} \subset \mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R} \subset \mathbb{C}$ .

## 3.2. Operationen

It is not the job of mathematicians to do correct arithmetical operations. It is the job of bank accountants.

---

*Samuil Shchatunovski*

Im vergangenen Kapitel haben wir einige Operationen definiert. Ganz allgemein sind Operationen mathematische Vorschriften, mit deren Hilfe man aus mathematischen Objekten neue Objekte erzeugen kann.

Operatoren werden häufig mit bestimmten Symbolen bezeichnet (für die Addition das Pluszeichen  $+$ , für die Multiplikation der Punkt  $\cdot$ , usw.). Meist handelt es sich um ein oder zwei Argumente, dann spricht man von ein- bzw. zweistelligen Operatoren.

Viele der von uns vorgestellten Operationen haben auch besondere Eigenschaften, weshalb wir hier gesondert auf einige Regeln und Gesetzmäßigkeiten eingehen wollen. Die meisten dieser Rechenvorschriften solltet ihr aus der Schule bereits kennen, wir führen sie hier aber der Vollständigkeit halber nochmals auf.

### 3.2.1. Potenz

Für das Rechnen mit Potenzen gelten einige Regeln, die sogenannten Potenzgesetze:

1.  $a^0 = 1$

Ein Sonderfall zu diesem Gesetz ist die Potenz  $0^0$ . Man kann argumentieren, dass  $0^x$  für alle  $x$  0 ist, da 0 beliebig oft mit sich selbst multipliziert 0 ist. Das widerspricht jedoch diesem Potenzgesetz, das besagt, dass  $x^0 = 1$  ist. In der gängigen Literatur wird daher häufig  $0^0$  gar nicht definiert. In einigen Anwendungsfällen, zum Beispiel bei bestimmten Grenzwerten oder Potenzreihen, ist es jedoch sinnvoll,  $0^0$  entweder als 0 oder als 1 zu definieren.

In der Informatik hat sich mittlerweile aus numerischen Gründen die Definition  $0^0 = 1$  durchgesetzt, sodass beispielsweise in Python `math.pow(0.0, 0.0) = 1.0` ist.

2.  $a^{b+c} = a^b \cdot a^c$

Dies ergibt sich aus  $a^{b+c} = \underbrace{a \cdot \dots \cdot a}_{b+c} = \underbrace{a \cdot \dots \cdot a}_b \cdot \underbrace{a \cdot \dots \cdot a}_c = a^b \cdot a^c$

3.  $a^{b-c} = \frac{a^b}{a^c}$

4.  $a^{-b} = \frac{1}{a^b}$

5.  $(a \cdot b)^c = a^c \cdot b^c$

Dies folgt mit der Kommutativität und Assoziativität der Multiplikation  $(a \cdot b)^c = \underbrace{(a \cdot b) \cdot \dots \cdot (a \cdot b)}_c = \underbrace{(a \cdot \dots \cdot a)}_c \cdot \underbrace{(b \cdot \dots \cdot b)}_c = a^c \cdot b^c$

### 3. Zahlenmengen und Operationen

$$6. \left(\frac{a}{b}\right)^c = \frac{a^c}{b^c}$$

$$7. (a^b)^c = a^{b \cdot c}$$

Dies kann man herleiten mittels:

$$(a^b)^c = \underbrace{(a^b) \cdot \dots \cdot (a^b)}_c = \underbrace{a \cdot \dots \cdot a}_b \cdot \dots \cdot \underbrace{a \cdot \dots \cdot a}_b = \underbrace{a \cdot \dots \cdot a}_{b \cdot c}$$

$$8. a^{\frac{c}{b}} = \sqrt[b]{a^c}$$

Um dies zu zeigen, muss man zuerst sehen, dass  $a^c = a^{\frac{c}{b} \cdot b} \stackrel{7}{=} (a^{\frac{c}{b}})^b$ . Mittels der Definition der Wurzel (Definition 3.7 auf Seite 22) ergibt sich daraus unsere gewünschte Formel.

#### 3.2.2. Wurzel

Mithilfe des oben eingeführten 8. Potenzgesetzes können alle Wurzelgesetze hergeleitet werden.

$$1. \sqrt[a]{b} \cdot \sqrt[a]{c} = \sqrt[a]{b \cdot c}$$

$$2. \frac{\sqrt[a]{b}}{\sqrt[a]{c}} = \sqrt[a]{\frac{b}{c}}$$

$$3. \sqrt[a]{\sqrt[b]{c}} = \sqrt[a \cdot b]{c}$$

$$4. (\sqrt[a]{b})^c = \sqrt[a]{b^c}$$

**Hinweis.** Man muss bei diesen Regeln aufpassen, wenn die Zahlen „unter der Wurzel“ negativ sind. Beispiel:

$$\sqrt{-1} \cdot \sqrt{-1} = i \cdot i = -1 \neq \sqrt{-1 \cdot -1} = 1$$

Die Wurzeln mit Exponent 2 und 3 nennt man Quadrat-, bzw. Kubikwurzel. Verkürzend kann man für die Quadratwurzel den Exponenten auch weglassen:

$$\sqrt{a} := \sqrt[2]{a}$$

#### 3.2.3. Logarithmus

$$1. \log_a(b \cdot c) = \log_a b + \log_a c$$

$$2. \log_a \frac{b}{c} = \log_a b - \log_a c$$

$$3. \log_a \frac{1}{b} = -\log_a b$$

$$4. \log_a b^c = c \log_a b$$

$$5. \log_a b = \frac{\ln b}{\ln a} = \frac{\log_x b}{\log_x a}$$

Spielt es keine Rolle, welche Basis der gerade verwendete Logarithmus hat (weil es aus dem Kontext ersichtlich ist oder nur die Größenordnung wichtig ist), so schreibt man auch nur  $\log a$ . Außerdem gibt es für einige zusätzliche, besonders häufig verwendete Logarithmen folgende Abkürzungen:

- $\ln a := \log_e a$  (logarithmus naturalis, natürlicher Logarithmus)
- $\lg a := \log_{10} a$  (dekadischer Logarithmus)
- $\text{ld } a := \log_2 a$  (logarithmus dualis)



## 4. Analysis

I could be bounded in a nutshell and count myself a king of infinite space.

---

*William Shakespeare*

Obwohl sich bereits die Mathe-Vorlesung im ersten Semester fast ausschließlich mit Analysis beschäftigen wird, gibt es viele Grundlagen, die im späteren Studium fehlen werden, wenn man sie nicht ausreichend vertieft. Daher wollen wir hier auf diese Grundlagen der Analysis nochmal eingehen. Dabei werden wir sowohl die Schulmathematik wiederholen, als auch, wo wir es für nötig halten, ein wenig darüber hinausgehen. Auf viele formale Beweise werden wir jedoch verzichten, diese werden in Mathematik I nachgeholt werden.

### 4.1. Funktionen

Einer der wichtigsten Begriffe der Mathematik ist der der *Funktion*. Ähnlich wie Relationen verknüpfen Funktionen zwei Mengen miteinander. Anders als eine Relation besitzt eine Funktion jedoch die Einschränkung, dass *jedem* Element aus der Definitionsmenge *genau ein* Wert aus der Wertemenge zugeordnet wird.

**Definition 4.1** (Funktion). Unter einer Funktion  $f$  versteht man eine Vorschrift, die jedem Element einer *Definitionsmenge*  $D$  genau ein Element einer *Wertemenge*  $W$  zuordnet. Man schreibt dann

$$f : D \rightarrow W, d \mapsto w$$

falls  $f$  das Element  $d \in D$  auf  $w \in W$  abbildet. Man schreibt dann auch  $f(d) = w$  und sagt, dass  $w$  und  $d$  *Bild* und *Urbild* voneinander unter (der Funktion)  $f$  sind.

Die Menge aller Elemente, in die die Funktion auch tatsächlich abbildet, nennt man *Bildmenge*. Diese Menge ist immer eine Teilmenge der Wertemenge, muss jedoch nicht mit ihr identisch sein.

**Hinweis.** Andere Begriffe für Definitionsmenge sind: Definitionsbereich, Domäne (engl. Domain). Andere Begriffe für Wertemenge sind: Wertebereich, Zielmenge.

**Lemma 4.2.** Eine Relation  $R \subseteq A \times B$  ist eine *Funktion*, falls es für jedes  $a \in A$  genau ein  $b \in B$  mit  $(a, b) \in R$  gibt.

Eine Funktion ist demnach nur ein Spezialfall einer *Relation* (Kapitel 2.2).

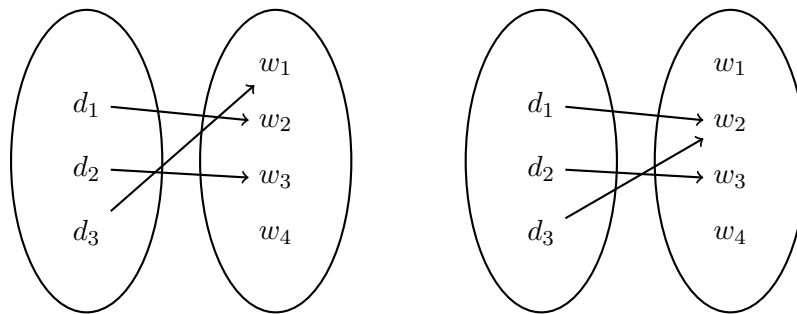
### 4.1.1. Eigenschaften von Funktionen

#### Injektivität, Surjektivität und Bijektivität

**Definition 4.3** (Injektivität). Eine Funktion  $f : D \rightarrow W$  heißt *injektiv*, falls für alle  $d_1, d_2 \in D$  gilt, dass aus  $f(d_1) = f(d_2)$  folgt, dass  $d_1 = d_2$ .

Intuitiv besagt die Eigenschaft also, dass für jedes Element  $w \in W$  höchstens ein  $d \in D$  existiert mit  $f(d) = w$ .

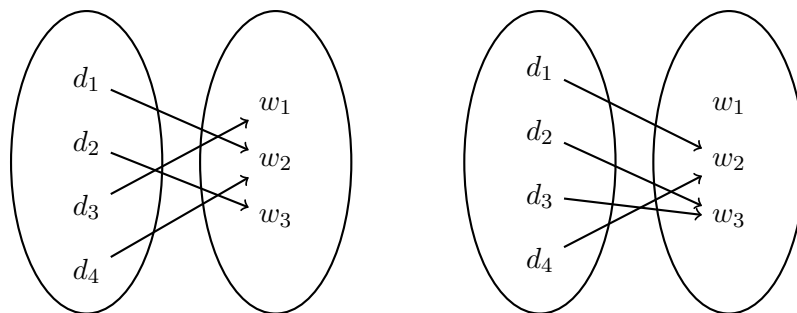
**Beispiel 4.4.** Die erste Funktion ist injektiv. Die zweite jedoch nicht, da das Element  $w_2$  die Injektivitätseigenschaft verletzt.



**Definition 4.5** (Surjektivität). Eine Funktion  $f : D \rightarrow W$  heißt *surjektiv*, falls für jedes  $w \in W$  mindestens ein  $d \in D$  existiert mit  $f(d) = w$ .

Intuitiv bedeutet Surjektivität also, dass zu jedem Element der Zielmenge mindestens ein Element im Definitionsbereich existiert.

**Beispiel 4.6.** Die erste Funktion ist surjektiv. Die zweite jedoch nicht, da es kein Element gibt, welches auf  $w_1$  abgebildet wird.



**Definition 4.7** (Bijektivität). Eine Funktion  $f : D \rightarrow W$  heißt *bijektiv*, falls  $f$  injektiv und surjektiv ist.

**Beispiel 4.8.** Die Funktion  $f : \mathbb{R} \rightarrow \mathbb{R}$  mit  $f(x) = x^2$  ist nicht injektiv. So ist z. B.  $f(2) = f(-2) = 4$ . Sie ist auch nicht surjektiv, da z. B.  $f(x) \neq -1$  für alle  $x \in \mathbb{R}$ . Die Funktion  $g : \mathbb{R} \rightarrow \mathbb{R}$  mit  $g(x) = 2x$  ist injektiv und surjektiv und damit auch bijektiv.

**Definition 4.9** (Umkehrfunktion). Sei  $f : D \rightarrow W$  eine bijektive Funktion, dann ist die Funktion  $f^{-1} : W \rightarrow D$  ihre *Umkehrfunktion*. Bei Umkehrfunktionen wird Bild und Urbild der ursprünglichen Funktion vertauscht. Der Funktionswert  $f^{-1}(y)$  ist also der Wert  $x \in D$ , für den  $f(x) = y$  gilt.

Aus dieser Definition folgt sofort, dass  $f^{-1}(f(x)) = x$  ist, da man das Bild von  $f$  sofort wieder in  $f^{-1}$  einsetzen kann und wieder  $x$  erhält.

### Beschränktheit, Monotonie und Stetigkeit

**Definition 4.10** (Beschränktheit). Eine Funktion  $f : D \rightarrow \mathbb{R}$  heißt *nach oben beschränkt*, falls es eine Zahl  $k \in \mathbb{R}$  gibt mit  $f(x) \leq k$  für alle  $x \in D$ .

$f$  heißt *nach unten beschränkt*, falls es eine Zahl  $b \in \mathbb{R}$  gibt mit  $b \leq f(x)$  für alle  $x \in D$ . Eine Funktion heißt *beschränkt*, falls  $f$  nach oben und nach unten beschränkt ist.

**Beispiel 4.11.** Die Funktion  $f : \mathbb{R} \rightarrow \mathbb{R}$  mit  $f(x) = \frac{1}{x^2}$  ist durch 0 nach unten beschränkt. Die Funktion  $g : \mathbb{R} \rightarrow \mathbb{R}$  mit  $g(x) = \sin x$  ist beschränkt (nach unten beschränkt durch  $-1$  und nach oben beschränkt durch  $1$ ).

**Definition 4.12** (Monotonie). Eine Funktion  $f : D \rightarrow \mathbb{R}$  heißt (*streng*) *monoton wachsend*, falls für alle  $x_1, x_2 \in D$  mit  $x_1 < x_2$  gilt, dass  $f(x_1) \leq f(x_2)$  (bzw.  $f(x_1) < f(x_2)$ ).

Eine Funktion  $f : D \rightarrow \mathbb{R}$  heißt (*streng*) *monoton fallend*, falls für alle  $x_1, x_2 \in D$  mit  $x_1 < x_2$  gilt, dass  $f(x_1) \geq f(x_2)$  (bzw.  $f(x_1) > f(x_2)$ ).

**Beispiel 4.13.** Die Funktion  $f : \mathbb{R} \rightarrow \mathbb{R}$  mit  $f(x) = 2x$  ist streng monoton wachsend. Die Funktion  $g : \mathbb{R} \rightarrow \mathbb{R}$  mit  $g(x) = 5$  ist monoton fallend, aber nicht *streng* monoton fallend.

**Definition 4.14** (Stetigkeit). Eine Funktion  $f : D \rightarrow \mathbb{R}$  heißt *stetig*, wenn eine beliebig kleine Änderung des Arguments nur eine beliebig kleine Änderung des Funktionswertes ergibt. Mit anderen Worten, es gibt keine „Sprünge“ im Funktionsgraph. Eine anschauliche Intuition ist, wenn man den Graph einer Funktion ohne Absetzen des Stiftes zeichnen kann.

### Symmetrie

**Definition 4.15** (Achsensymmetrie). Eine Funktion  $f$  heißt *achsensymmetrisch* zur y-Achse, falls für alle  $x \in D$  gilt, dass  $f(x) = f(-x)$ , d.h. die Funktion entspricht ihrer Spiegelung an der y-Achse.

**Hinweis.** Da Funktionen rechtseindeutig sind, also keinem x-Wert mehrere y-Werte zugeordnet werden, gibt es keine x-achsensymmetrische Funktionen. Jedoch kann es z.B. Symmetrieachsen geben, die parallel zur y-Achse liegen.

**Definition 4.16** (Punktsymmetrie). Eine Funktion  $f$  heißt *punktsymmetrisch* zum Ursprung, falls für alle  $x \in D$  gilt, dass  $f(-x) = -f(x)$ .

#### 4. Analysis

**Hinweis.** Andere Punktsymmetrien können z.B. mittels Verschiebung des vermuteten Symmetriepunkts auf den Ursprung überprüft werden, sind aber seltener interessant.

**Beispiel 4.17** (Symmetrie von Funktionen). Die Funktion  $\cos : \mathbb{R} \rightarrow \mathbb{R}$  ist *achsensymmetrisch*, aber nicht *punktsymmetrisch*.

Die Funktion  $\sin : \mathbb{R} \rightarrow \mathbb{R}$  ist *punktsymmetrisch*, aber nicht *achsensymmetrisch*.

##### 4.1.2. Transformationen

Immer wieder kommt es vor, dass man von einer Funktion nur die Formel und nicht den Funktionsgraphen sieht. Der Funktionsgraph einer Funktion ist die Menge aller Paare  $(x, f(x))$ . An der Uni könnt Ihr in der Regel keinen graphischen Taschenrechner verwenden. Um grob abzuschätzen, wie eine Funktion aussieht, ist es deshalb hilfreich zu wissen, wie sich Veränderungen an Teilen des Funktionsterms auf den Graph auswirken. Unter diese sogenannten Transformationen fallen:

- Streckung/Stauchung in x-Richtung um den Faktor  $a$ .

$$f^*(x) = f\left(\frac{x}{a}\right)$$

- Spiegelung an der y-Achse

$$f^*(x) = f(-x)$$

- Streckung/Stauchung in y-Richtung um den Faktor  $a$ .

$$f^*(x) = a \cdot f(x)$$

- Spiegelung an der x-Achse

$$f^*(x) = -f(x)$$

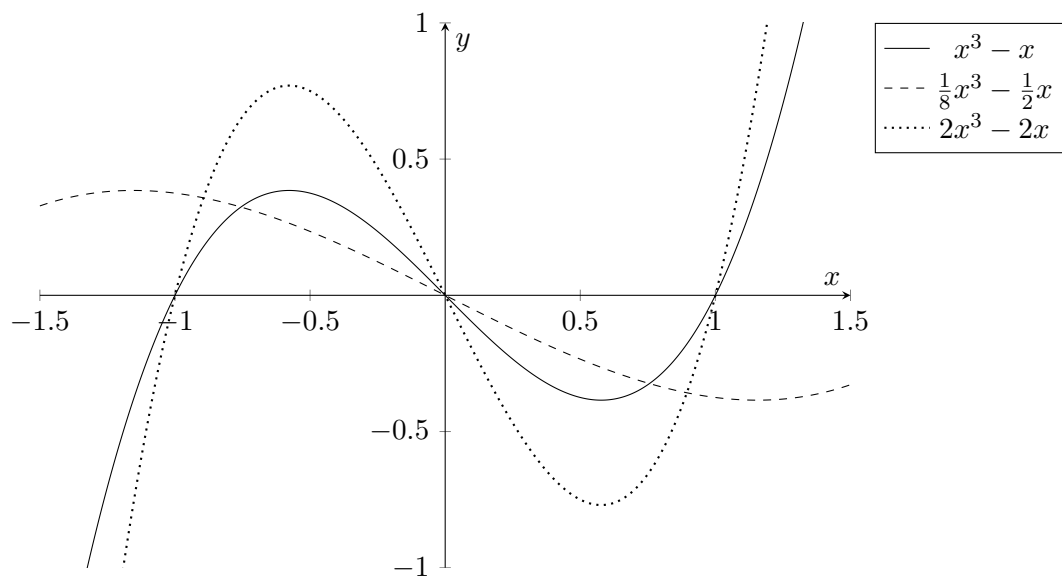
- Verschiebung/Translation in x-Richtung um die Konstante  $k$

$$f^*(x) = f(x - k)$$

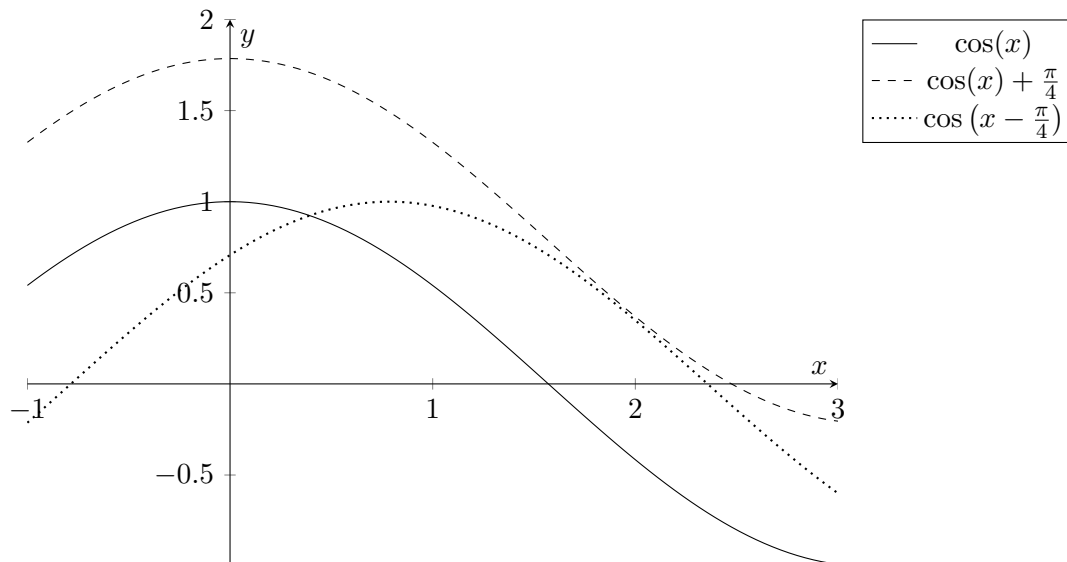
Zu beachten ist hierbei, dass, wenn man die Funktion um  $k$  in positiver Richtung verschieben will,  $k$  von  $x$  subtrahiert werden muss.

- Verschiebung in y-Richtung um die Konstante  $k$

$$f^*(x) = f(x) + k$$



**Abbildung 4.1.:** Die Funktion  $x^3 - x$  jeweils um den Faktor 2 in x- bzw. y-Richtung gestreckt.



**Abbildung 4.2.:** Die Funktion  $\cos(x)$  jeweils um die Konstante  $\frac{\pi}{4}$  in x- bzw. y-Richtung verschoben.

## 4.2. Folgen

Funktionen mit Definitionsmenge  $\mathbb{N}$  bilden aufgrund der Abzählbarkeit der natürlichen Zahlen eine Auflistung ihrer Funktionswerte. Da man also eindeutig sagen kann, welcher Funktionswert der „erste“, „zweite“, usw. ist, spricht man auch von einer *Folge*.

**Definition 4.18** (Folge). Eine Folge ist eine Abbildung  $a : \mathbb{N} \rightarrow W$ .

Von besonderer Bedeutung sind für uns Folgen der Form  $a : \mathbb{N} \rightarrow \mathbb{R}$  bzw.  $a : \mathbb{N} \rightarrow \mathbb{C}$  (Zahlenfolgen).

Man schreibt für die Folge meist  $(a_n)_{n \in \mathbb{N}}$ . Der Wert an der Stelle  $n$  heißt *n-tes Folgenglied*, anstatt  $a(n)$  schreibt man bei Folgen meist  $a_n$ .

**Beispiel 4.19.** Die Folge  $a_n = \frac{1}{n}$  für  $n \neq 0$  ist in Abbildung 4.3 visualisiert.

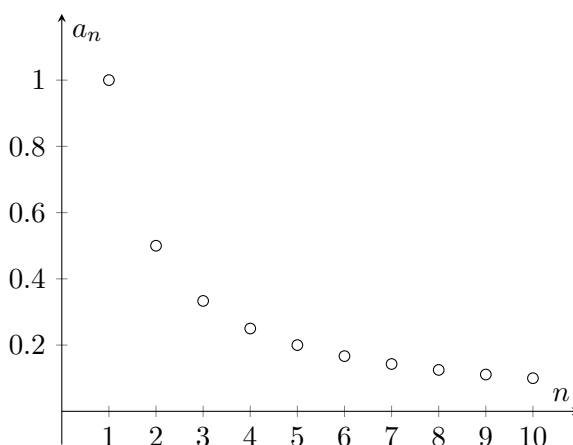


Abbildung 4.3.: Die Folge  $a_n = \frac{1}{n}$ .

**Hinweis.** Es gibt auch Datenbanken, in denen Folgen abgerufen werden können. Eine bekannte ist The On-Line Encyclopedia of Integer Sequences<sup>TM1</sup>.

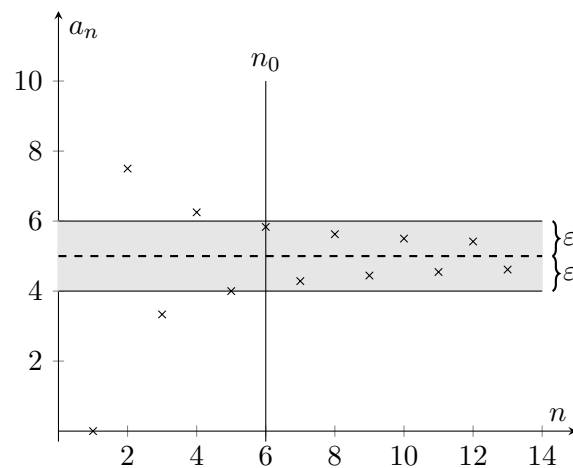
### 4.2.1. Konvergenz

Der Begriff der *Konvergenz* ist einer der wichtigsten der Analysis. Intuitiv gesagt bedeutet Konvergenz, dass sich eine Folge für größer werdende  $n$  immer mehr einem gewissen Wert annähert. Mit anderen Worten: es gibt eine Zahl  $n_0$ , ab der alle nachfolgenden Folgenglieder nur noch eine Differenz von höchstens  $\varepsilon > 0$  zum Grenzwert haben, wobei  $\varepsilon$  beliebig klein werden kann.

Man kann den Grenzwert auch so verstehen, dass für jedes  $\varepsilon$  nur endlich viele Folgenglieder einen Abstand größer  $\varepsilon$  zum Grenzwert haben, also außerhalb eines sogenannten *Epsilonschlauchs* liegen, siehe Abbildung 4.4.

Wir wollen die Konvergenz jedoch auch formal definieren:

<sup>1</sup><http://oeis.org/>



**Abbildung 4.4.:** Der Grenzwert der Folge  $a_n = \frac{(-1)^n \cdot 5}{n} + 5$  ist offenbar 5. Für  $\varepsilon = 1$  ist  $n_0 = 6$ , da ab diesem Wert alle folgenden Glieder innerhalb des Epsilonschlauchs liegen.

**Definition 4.20** (Konvergenz einer Folge). Sei  $(a_n)_{n \in \mathbb{N}}$  eine Folge, die in die reellen Zahlen abbildet und  $a \in \mathbb{R}$ . Dann ist

$$\lim_{n \rightarrow \infty} a_n = a$$

der *Grenzwert* der Folge, falls für *jedes*  $\varepsilon \in \mathbb{R}$  mit  $\varepsilon > 0$  ein  $n_0 \in \mathbb{N}$  existiert, sodass für alle  $n \in \mathbb{N}$  mit  $n \geq n_0$  gilt, dass  $|a_n - a| < \varepsilon$ .

Existiert keine solche Zahl  $a \in \mathbb{R}$ , dann heißt die Folge *divergent*.

**Beispiel 4.21** (Konvergenz einer Folge). Der Grenzwert der Folge  $a_n = \frac{1}{n}$  ist 0.

*Beweis.* Sei  $\varepsilon > 0$ . Wähle nun  $n_0 \in \mathbb{N}$  sodass  $n_0 > \frac{1}{\varepsilon}$ . Diese Zahl existiert für jedes  $\varepsilon > 0$ :  $\frac{1}{\varepsilon}$  wird für kleiner werdende  $\varepsilon$  immer größer, ist jedoch stets eine reelle Zahl. Man wählt also einfach die nächste natürliche Zahl, die größer ist.

Dann ist für  $n \geq n_0$ :

$$\begin{aligned} \left| \frac{1}{n} - 0 \right| &= \frac{1}{n} \\ &\leq \frac{1}{n_0} \\ &< \varepsilon \end{aligned}$$

□

### 4.2.2. Wichtige Folgen

Es gibt einige bestimmte Formen von Folgen, die von besonderer Bedeutung sind. Auf diese wollen wir nun im Speziellen eingehen.

## 4. Analysis

### Arithmetische Folge

Eine Folge, bei der alle Folgenglieder die gleiche Differenz zu ihrem Nachfolger haben, heißt Arithmetische Folge. Eine solche Folge ist beispielsweise  $a_n = 2n$ , also die Folge der geraden Zahlen. Hier hat jedes Folgenglied die Differenz 2 zu seinem Nachfolger. Ganz allgemein haben arithmetische Folgen die Form

$$a_n = a_0 + d \cdot n$$

Die Differenz zwischen den einzelnen Folgengliedern kann man dann direkt am Parameter  $d$  ablesen. Da  $d \cdot 0 = 0$ , ist  $a_0$  hier direkt das erste Folgenglied<sup>2</sup>.

### Potenzfolge

Folgen der Form

$$a_n = n^k$$

nennt man Potenzfolgen. Für  $k = 1$  entsteht die arithmetische Folge  $a_n = n$ . Potenzfolgen können umgeschrieben werden als  $a_n = \frac{1}{n^{-k}}$ . Für negative  $k$  konvergieren sie gegen 0. Für  $k = -1$  ergibt sich die *Harmonische Folge*  $a_n = \frac{1}{n^k}$ .

### Geometrische Folgen

Bei arithmetischen Folgen haben die Folgenglieder untereinander immer die gleiche Differenz. Geometrische Folgen sind dadurch charakterisiert, dass zwei aufeinanderfolgende Folgenglieder immer das gleiche Verhältnis zueinander haben. Anders ausgedrückt, der Quotient  $\frac{a_{n+1}}{a_n} = q$  ist für alle  $n$  immer der gleiche.

Formal ist eine geometrische Folge definiert durch

$$a_n = a_0 \cdot q^n$$

Auch hier bezeichnet  $a_0$  direkt das erste Folgenglied, da  $q^0 = 1$  ist und man somit  $a_0$  erhält. Geometrische Folgen konvergieren für  $|q| < 1$  gegen 0, für  $q = 1$  gegen  $a_0$ , und für alle anderen Werte von  $q$  divergieren sie.

---

<sup>2</sup>Hier sieht man, warum es mitunter ungünstig ist,  $0 \in \mathbb{N}$  zu definieren: es ist unintuitiv Abzählungen bei 0 zu beginnen.  $a_0$  kann erstes oder „nullstes“ Element der Folge genannt werden.

### 4.2.3. Rechenregeln für Zahlenfolgen

**Satz 4.22.** Seien  $(a)_{n \in \mathbb{N}}, (b)_{n \in \mathbb{N}}$  konvergente Zahlenfolgen mit  $\lim_{n \rightarrow \infty} a_n = a$  und  $\lim_{n \rightarrow \infty} b_n = b$  und  $c \in \mathbb{R}$ . Dann:

$$\begin{aligned}\lim_{n \rightarrow \infty} (a_n \pm b_n) &= a \pm b \\ \lim_{n \rightarrow \infty} (a_n \cdot b_n) &= a \cdot b \\ \lim_{n \rightarrow \infty} \sqrt{a_n} &= \sqrt{a} \text{ für } a \geq 0 \\ \lim_{n \rightarrow \infty} c \cdot a_n &= c \cdot a \\ \lim_{n \rightarrow \infty} \frac{1}{a_n} &= \frac{1}{a} \text{ für } a \neq 0\end{aligned}$$

*Beweis.* Wir beweisen beispielhaft  $\lim_{n \rightarrow \infty} (a_n + b_n) = a + b$ .

Da  $\lim_{n \rightarrow \infty} a_n = a$  gilt, gibt es für jedes  $\varepsilon > 0$  ein  $n_0$ , sodass für alle  $n \in \mathbb{N}$  mit  $n \geq n_0$  gilt  $|a_n - a| < \varepsilon$ . Da  $\lim_{n \rightarrow \infty} b_n = b$  gilt, gibt es für jedes  $\varepsilon > 0$  ein  $m_0$ , sodass für alle  $n \in \mathbb{N}$  mit  $n \geq m_0$  gilt  $|b_n - b| < \varepsilon$ .

Sei nun  $\varepsilon > 0$ . Dann ist auch  $\frac{\varepsilon}{2} > 0$ . Dann gibt es für  $\frac{\varepsilon}{2}$  ein  $n'_0$ , sodass  $|a - a_n| < \frac{\varepsilon}{2}$  für  $n \geq n'_0$  und ein  $n''_0$ , sodass  $|b - b_n| < \frac{\varepsilon}{2}$  für  $n \geq n''_0$ . Wähle nun  $n_0 := \max\{n'_0, n''_0\}$ . Sei  $n \geq n_0$ :

$$\begin{aligned}|a_n + b_n - a - b| &= |a_n - a + b_n - b| \\ &\leq |a_n - a| + |b_n - b| \\ &< \frac{\varepsilon}{2} + \frac{\varepsilon}{2} \\ &= \varepsilon\end{aligned}$$

□

### 4.2.4. Rekursion

In order to understand recursion, one must first understand recursion.

---

Andrew Koenig

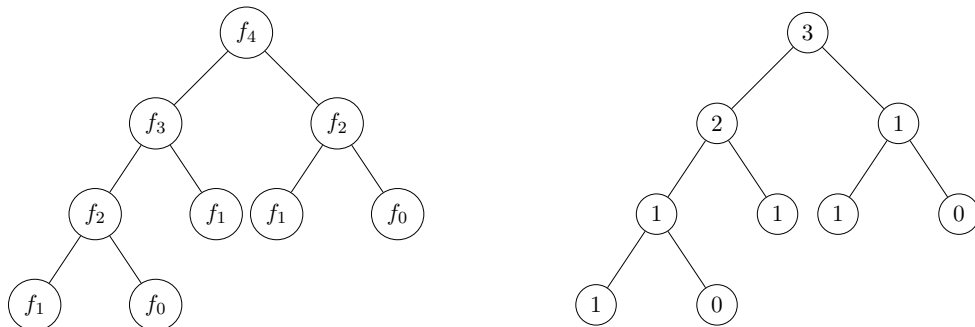
Bisher haben wir immer Folgen gesehen, bei denen man das  $n$ -te Glied direkt durch Einsetzen in eine Rechenvorschrift berechnen kann. Man kann aber auch Folgen definieren, die für die Berechnung des  $n$ -ten Gliedes auf vorherige Glieder zurückgreifen. Dieses Vorgehen nennt man *Rekursion*, solche Folgen heißen *rekursive Folgen*.

#### 4. Analysis

**Beispiel 4.23** (Fibonacci-Folge). Eine sehr bekannte rekursive Folge ist die Fibonacci-Folge  $(f_n)_{n \in \mathbb{N}}$  (benannt nach *Leonardo Fibonacci*<sup>3</sup>). Sie entsteht durch Addition der jeweils beiden vorhergehenden Folgenglieder. Formal ist die Folge folgendermaßen definiert:

$$\begin{aligned} f_0 &= 0 \\ f_1 &= 1 \\ f_n &= f_{n-1} + f_{n-2} \text{ für } n \geq 2 \end{aligned}$$

Die ersten Folgenglieder lauten  $(0, 1, 1, 2, 3, 5, 8, 13, 21, 34, \dots)$ , wie man sie berechnet ist in dem folgenden Baum kurz graphisch dargestellt.



Die Fibonacci-Folge taucht an vielen Stellen in der Natur auf. So sind beispielsweise bei vielen Pflanzen die Samen in den Blütenständen (zum Beispiel die der Sonnenblume) in einer Spirale angeordnet, deren Armlängen Fibonacci-Zahlen sind.

Hergeleitet hat Leonardo Fibonacci die Folge durch das Wachstum einer Kaninchen-Population. Er nahm dazu an, dass jedes erwachsene Kaninchenpaar ein Junges pro Monat produziert. Neugeborene Kaninchen brauchen einen Monat bis zur Geschlechtsreife. Bei unsterblichen Kaninchen ergibt sich dann die Anzahl der Kaninchen in Abhängigkeit der Monate genau durch die Fibonacci-Folge.

**Beispiel 4.24** (Fakultät). Die *Fakultät* (geschrieben  $n!$ ) ist folgendermaßen definiert:

$$\begin{aligned} 0! &= 1 \\ n! &= n \cdot (n-1)! \text{ für } n \geq 1 \end{aligned}$$

Einige Werte der Fakultät:

$$\begin{aligned} 0! &= 1 \\ 1! &= 0! \cdot 1 = 1 \cdot 1 = 1 \\ 2! &= 1! \cdot 2 = 1 \cdot 1 \cdot 2 = 2 \\ 3! &= 2! \cdot 3 = 1 \cdot 1 \cdot 2 \cdot 3 = 6 \\ 11! &= 1 \cdot 1 \cdot 2 \cdot 3 \cdot 4 \cdot 5 \cdot 6 \cdot 7 \cdot 8 \cdot 9 \cdot 10 \cdot 11 = 39,916,800 \end{aligned}$$

<sup>3</sup>italienischer Mathematiker, ca. 1175 - ca. 1240

### 4.3. Reihen

Als Motivation für Reihen wollen wir das Problem von Achilles und der Schildkröte betrachten, welches vom griechischen Philosophen *Zenon von Elea*<sup>4</sup> erdacht wurde.

Achilles läuft ein Wettrennen gegen eine Schildkröte und gibt dieser einen Meter Vorsprung. Zenon behauptete nun, dass Achilles das Wettrennen nie gewinnen könne, egal wie schnell er laufen würde. Er argumentierte dabei wie folgt:

Wenn Achilles  $k$ -mal so schnell laufen kann wie die Schildkröte, dann erreicht Achilles die Startposition der Schildkröte wenn diese bereits  $\frac{1}{k}$  Meter weitergekröchen ist. Erreicht Achilles nun die neue Position der Schildkröte, so ist diese bereits um  $\frac{1}{k^2}$  Meter weitergelaufen. Diesen Weg muss Achilles nun wieder aufholen usw.

Die Strecke, die Achilles aufholen müsste, würde sich dann aufsummieren zu

$$1 + \frac{1}{k} + \frac{1}{k^2} + \frac{1}{k^3} + \dots$$

Dieser Weg würde sich ewig weiter aufsummieren, damit kann Achilles die Schildkröte nie einholen.

Der Denkfehler des Philosophen ist, dass er davon ausging, dass die Summe von unendlich vielen positiven Zahlen selbst nie eine endliche Zahl sein kann. Mit der Definition des Grenzwertes können wir dieses Problem jedoch formal auffassen und werden sehen, dass die obige Summe, die aus unendlich vielen Gliedern besteht, endlich ist.

#### 4.3.1. Summen und Produkte

Bevor wir uns um Reihen kümmern, müssen wir jedoch eine neue mathematische Notation einführen. Denn oft will man alle Glieder einer Zahlenfolge bis zu einem gewissen Index aufsummieren. Sei  $(a_n)_{n \in \mathbb{N}}$  eine Folge. Will man alle Folgenglieder der Folge von  $k$  bis  $n$  (für  $k, n \in \mathbb{N}$ ) aufsummieren, so verwendet man häufig das *Summenzeichen*  $\Sigma$ .

**Definition 4.25** (Summenzeichen). Für eine Zahlenfolge  $(a_n)_{n \in \mathbb{N}}$  definiert man

$$\sum_{i=k}^n a_i := a_k + a_{k+1} + \dots + a_n$$

Ist  $k > n$ , so ist  $\sum_{i=k}^n a_i = 0$ .

#### Rechenregeln für Summen

Summen haben einige nützliche Eigenschaften, welche helfen mit Summen zu rechnen und verhindern, dass man mühsam jedes Glied von Hand addieren muss.

- Da die Addition assoziativ ist, kann eine Summe auch in zwei Summen zerteilt und unabhängig voneinander berechnet werden.

$$\sum_{i=k}^n a_i = \sum_{i=k}^l a_i + \sum_{i=l+1}^n a_i \text{ für } k \leq l \leq n$$

---

<sup>4</sup>5. Jh. v. Chr.

#### 4. Analysis

- Ein Spezialfall ist hier, dass auch nur das erste oder letzte Glied herausgezogen werden kann.

$$\sum_{i=k}^{n+1} a_i = a_{n+1} + \sum_{i=k}^n a_i$$

- Nach dem Distributivgesetz kann eine multiplikative Konstante aus der Summe ausgeklammert werden.

$$\sum_{i=k}^n (c \cdot a_i) = c \cdot \sum_{i=k}^n a_i$$

- Ebenso kann man additive Konstanten aus der Summe herausziehen. Da diese pro Summenglied je einmal vorkommen, müssen sie mit der Anzahl der Summenglieder multipliziert werden.

$$\sum_{i=k}^n (c + a_i) = (n - k + 1) \cdot c + \sum_{i=k}^n a_i$$

- Manchmal kann es sogar nützlich sein, den Index der Summe zu verschieben. Dann muss allerdings auch der innere Teil der Summe passend modifiziert werden.

$$\sum_{i=k}^n a_i = \sum_{i=k+m}^{n+m} a_{i-m}, \text{ für } m \in \mathbb{Z}$$

**Hinweis.** Auch für die Multiplikation von Folgengliedern gibt es eine verkürzende Schreibweise. In diesem Fall verwendet man das Zeichen  $\prod$ :

$$\prod_{i=k}^n a_i := a_k \cdot a_{k+1} \cdot \dots \cdot a_n$$

Ein Spezialfall ist das leere Produkt  $\prod_{i=k}^n a_i = 1$  für  $n < k$ .

**Beispiel 4.26** (Kleiner Gauß). Es gilt

$$\sum_{i=1}^n i = \frac{n \cdot (n+1)}{2}$$

Diese Aussage werden wir später beweisen, siehe Satz 6.23.

##### 4.3.2. Reihen

Reihen sind nichts anderes als Summen über Folgenglieder. Da wir beide Begriffe schon eingeführt haben, können wir sie nun verwenden, um eine ganz einfache Definition von Reihen anzugeben.

**Definition 4.27** (Reihe). Es sei  $(a_n)_{n \in \mathbb{N}}$  eine Folge. Dann lässt sich eine neue Folge  $(S_n)_{n \in \mathbb{N}}$  wie folgt definieren

$$S_n := \sum_{i=0}^n a_i$$

Die Glieder  $S_n$  nennt man auch *Partialsummen* und die Folge der Partialsummen nennt man *Reihe*. Wenn die Folge  $(S_n)_{n \in \mathbb{N}}$  konvergiert (divergiert), dann sagt man, dass die Reihe konvergiert (divergiert). Man definiert die Summe oder den Wert der Reihe  $S$  als

$$S := \lim_{n \rightarrow \infty} S_n = \lim_{n \rightarrow \infty} \sum_{i=0}^n a_i =: \sum_{i=0}^{\infty} a_i$$

### 4.3.3. Wichtige Reihen

#### Arithmetische Reihe

Sei  $(a_n)_{n \in \mathbb{N}}$  eine arithmetische Folge (Kapitel 4.2.2). Dann wird  $(S_n)_{n \in \mathbb{N}}$  als arithmetische Reihe bezeichnet:

$$S_n = \sum_{i=0}^n a_0 + c \cdot i$$

mit  $c \in \mathbb{R}$ . Die Reihe ist im Allgemeinen divergent.

#### Potenzreihe

Eine Potenzreihe ist eine Reihe  $(S_n)_{n \in \mathbb{N}}$  der Form:

$$S_n = \sum_{i=0}^n a_n \cdot (x - x_0)^i$$

Hierbei kann  $(a_n)_{n \in \mathbb{N}}$  eine beliebige Folge sein.

Mit den Konvergenzeigenschaften dieser wunderschönen Reihe werdet ihr euch noch in den Mathematik-Vorlesungen eingehend beschäftigen.

#### Geometrische Reihe

Sei  $(a_n)_{n \in \mathbb{N}}$  eine geometrische Folge (Kapitel 4.2.2). Dann wird  $(S_n)_{n \in \mathbb{N}}$  als geometrische Reihe bezeichnet:

$$S_n = \sum_{i=0}^n a_0 \cdot q^i$$

und der Wert der Reihe ist für  $|q| < 1$

$$S = \frac{a_0}{1 - q}$$

## 4. Analysis

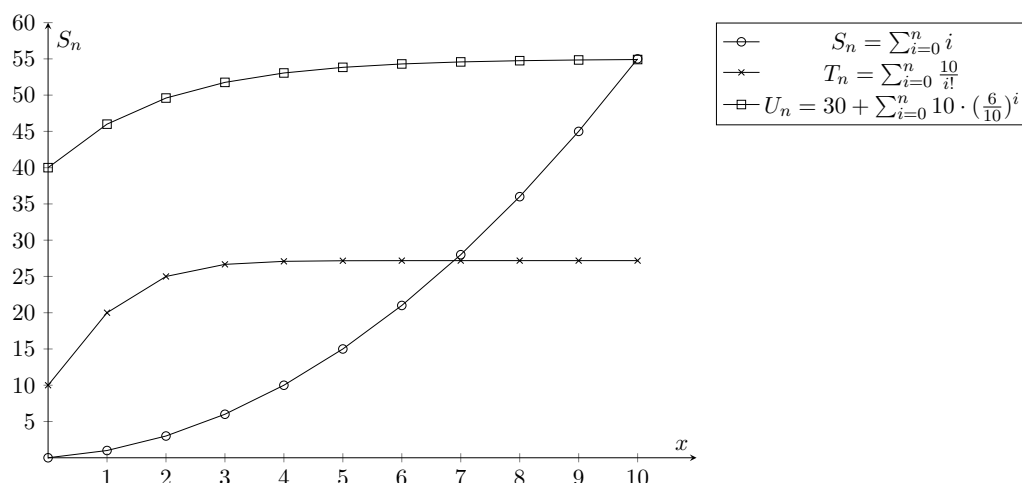


Abbildung 4.5.: Graphische Darstellung einiger Reihen.

Damit löst sich auch das scheinbare Paradoxon um Achilles und die Schildkröte. Ist  $k > 1$  (da Achilles ja schneller läuft als die Schildkröte), so ist der Weg, den Achilles aufholen muss, gegeben durch eine Geometrische Reihe mit  $q = \frac{1}{k}$  und somit  $q < 1$  und  $a_0 = 1$ , dieser Wert ist endlich, Achilles kann also gewinnen.

**Definition 4.28** (Eulersche Zahl). Eine der wichtigsten Konstanten der Mathematik ist die *Eulersche Zahl*  $e$ . Sie findet in vielen natürlichen Wachstumsprozessen Anwendung, ist die Basis des natürlichen Logarithmus und wird auch zur Darstellung imaginärer Zahlen benutzt. Man definiert sie über den Grenzwert einer Reihe:

$$e := \sum_{i=0}^{\infty} \frac{1}{i!}$$

Der zehnfache Wert dieser Reihe ist in Abbildung 4.5 mit Kreuzen dargestellt. Ihr Wert ist  $e \approx 2.71828 \dots$

## 4.4. Funktionsklassen

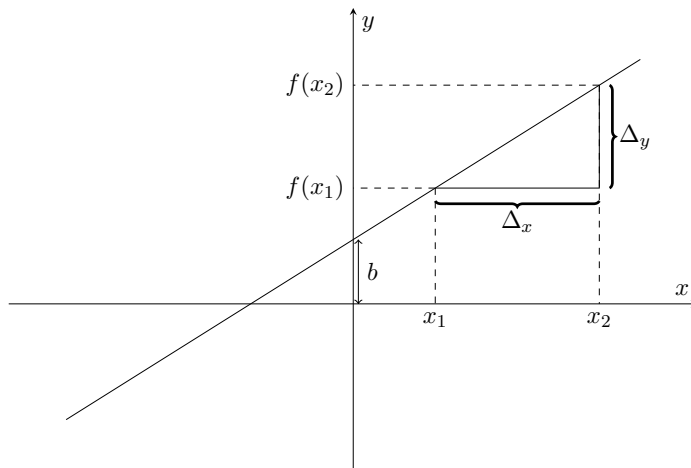
Von besonderem Interesse sind oft Funktionen, die von einer Zahlenmenge in eine andere Zahlenmenge abbilden (besonders  $\mathbb{R}, \mathbb{C}$ ). Im Folgenden wollen wir einige Klassen solcher Funktionen betrachten und auf einige Eigenschaften eingehen.

### 4.4.1. Lineare Funktionen

Eine lineare Funktion  $f$  ist eine Funktion der Form

$$f(x) = m \cdot x + b$$

mit  $m, b \in \mathbb{R}$ . Der Funktionsgraph einer linearen Funktion ist in Abbildung 4.6 dargestellt.



**Abbildung 4.6.:** Funktionsgraph einer linearen Funktion  $f(x) = mx + b$ .

Man bezeichnet dabei  $b$  als *y-Achsenabschnitt* und  $m$  als *Steigung* der Funktion. Der Punkt  $(0, b)$  gibt stets den Schnittpunkt der Funktion mit der y-Achse an.

Hat man von einer linearen Funktion zwei Punkte  $f(x_1) = y_1$  und  $f(x_2) = y_2$  gegeben, dann kann man die Steigung einfach aus dem Quotienten  $m = \frac{\Delta y}{\Delta x}$  berechnen:

$$\begin{aligned} \frac{\Delta y}{\Delta x} &= \frac{f(x_2) - f(x_1)}{x_2 - x_1} \\ &= \frac{m \cdot x_2 + b - m \cdot x_1 - b}{x_2 - x_1} \\ &= \frac{m \cdot (x_2 - x_1)}{x_2 - x_1} \\ &= m \end{aligned}$$

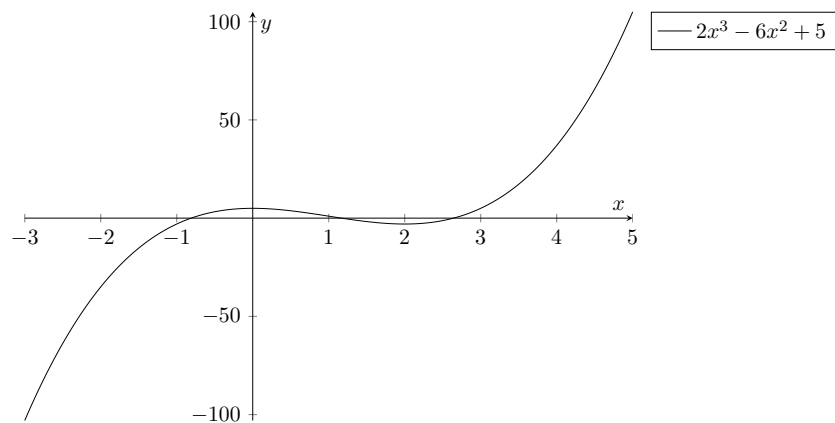
#### 4.4.2. Polynome

Ein Polynom  $p$  vom Grad  $n$  ist eine Funktion der Form

$$p(x) = \sum_{i=0}^n a_i \cdot x^i$$

mit  $a_i \in \mathbb{R}$  für  $0 \leq i \leq n$  und  $a_n \neq 0$ . Eine lineare Funktion ist somit nur ein Spezialfall eines Polynoms. Polynome spielen in vielen Anwendungsbereichen eine wichtige Rolle (z. B. das sogenannte *Taylor-Polynom*, Laufzeiten von Algorithmen). Ein Polynom vom Grad  $n$  hat in  $\mathbb{R}$  *höchstens*  $n$  Nullstellen, also Stellen  $x' \in \mathbb{R}$  für die gilt  $f(x') = 0$ . Der Funktionsgraph eines Polynoms ist in Abbildung 4.7 exemplarisch dargestellt.

#### 4. Analysis



**Abbildung 4.7.:** Funktionsgraph des Polynoms  $p(x) = 2x^3 - 6x^2 + 5$ .  $p$  ist ein Polynom vom Grad 3.

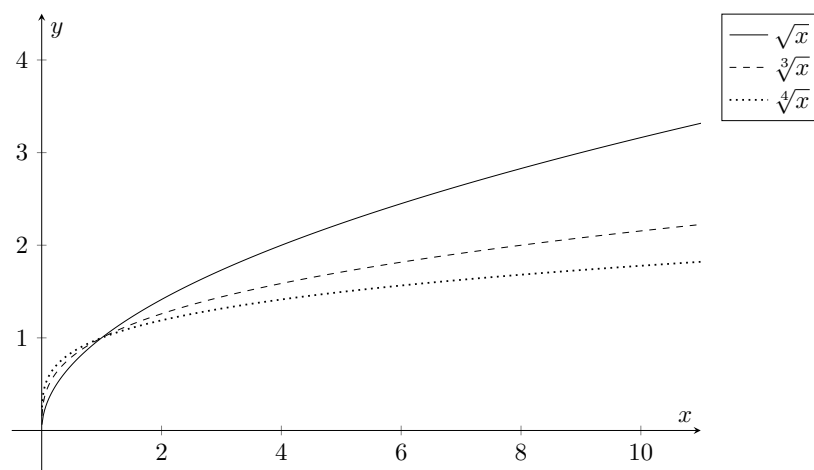
#### 4.4.3. Wurzelfunktionen

Wie wir bereits in Definition 3.7 auf Seite 22 gesehen haben, ist die  $n$ -te Wurzel von  $x$  definiert durch

$$\sqrt[n]{x} = a \text{ gdw. } a^n = x$$

Um eine Funktion zu erhalten, müssen wir sicherstellen, dass diese Gleichung *genau eine* Lösung hat. Wenn  $n$  eine gerade Zahl ist, so ist aber stets  $a$  als auch  $-a$  eine Lösung. Die Wurzelfunktion bildet deshalb nur auf die positive Lösung ab und ist daher eine Funktion von  $\mathbb{R}_0^+$  nach  $\mathbb{R}_0^+$ .

Drei Graphen von Wurzelfunktionen sind in Abbildung 4.8 dargestellt.



**Abbildung 4.8.:** Funktionsgraphen von drei Wurzelfunktionen.

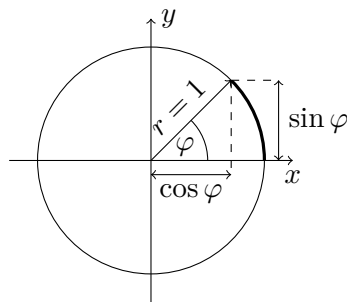
| $x$       | 0  | $\frac{\pi}{6}$      | $\frac{\pi}{4}$      | $\frac{\pi}{3}$      | $\frac{\pi}{2}$ |
|-----------|----|----------------------|----------------------|----------------------|-----------------|
|           | 0° | 30°                  | 45°                  | 60°                  | 90°             |
| $\sin(x)$ | 0  | $\frac{1}{2}$        | $\frac{1}{\sqrt{2}}$ | $\frac{\sqrt{3}}{2}$ | 1               |
| $\cos(x)$ | 1  | $\frac{\sqrt{3}}{2}$ | $\frac{1}{\sqrt{2}}$ | $\frac{1}{2}$        | 0               |
| $\tan(x)$ | 0  | $\frac{1}{\sqrt{3}}$ | 1                    | $\sqrt{3}$           | $\infty$        |

Tabelle 4.1.: Einige Funktionswerte der cos- und sin- Funktionen.

#### 4.4.4. Trigonometrische Funktionen

Trigonometrische Funktionen beschreiben geometrische Zusammenhänge in Dreiecken. Grundlage für die bekannten Funktionen cos und sin (Kosinus und Sinus) ist der *Einheitskreis* (Kreis mit Radius 1 um den Punkt  $(0,0)$  und Umfang  $2\pi$ ).

Stellen wir uns nun einen Zeiger vor, der vom Punkt  $(0,0)$  auf den Punkt  $(1,0)$  zeigt. Diesen Zeiger können wir um einen *Winkel*  $0 \leq |\varphi| \leq 2\pi$  drehen ( $\varphi \geq 0$  gegen den Uhrzeigersinn und  $\varphi < 0$  für Drehung *im* Uhrzeigersinn). Dabei ist  $\varphi$  die Länge des entsprechenden Bogens. Die  $x$ -Koordinate des neuen Punktes bezeichnet man mit  $\cos \varphi$  und die  $y$ -Koordinate mit  $\sin \varphi$ , siehe Abbildung 4.9.

Abbildung 4.9.: Drehung des Zeigers um  $\varphi = \frac{\pi}{4} \hat{=} 45^\circ$ .

Wir können diese Funktionen erweitern auf beliebige Zahlen  $x \in \mathbb{R}$  indem wir nach einer Umrundung des Kreises einfach weiter in die entsprechende Richtung laufen. Insgesamt erhalten wir also zwei Funktionen  $\cos : \mathbb{R} \rightarrow \mathbb{R}$  und  $\sin : \mathbb{R} \rightarrow \mathbb{R}$ .

Eine weitere Funktion, den Tangens, erhalten wir durch einfaches dividieren von Sinus und Kosinus:

$$\tan x := \frac{\sin x}{\cos x}$$

für  $\cos x \neq 0$ , also  $x \neq \pi n \pm \frac{\pi}{2}$ ,  $n \in \mathbb{N}$ .

Die Funktionsgraphen sind in Abbildung 4.10 dargestellt. Einige wichtige Funktionswerte finden sich in Tabelle 4.1.

## 4. Analysis

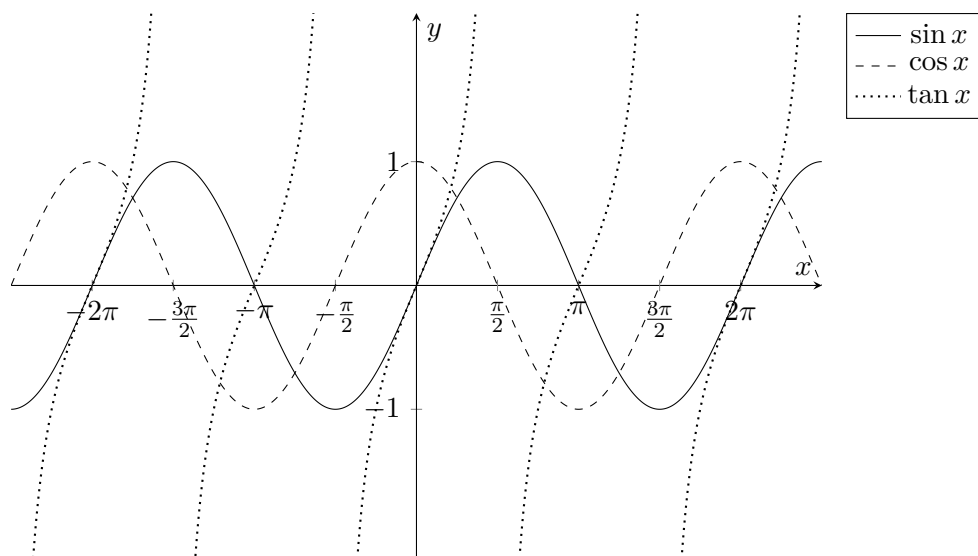


Abbildung 4.10.: Funktionsgraph der Sinus-, Kosinus- und Tangens-Funktion.

### 4.4.5. Exponentialfunktionen

Die *Eulersche Zahl*  $e$  haben wir bereits vorher definiert (Definition 4.28 auf Seite 44). Die *Exponentialfunktion*  $e : x \mapsto e^x$  ist für alle  $x \in \mathbb{R}$  definiert. Auch hier möchten wir auf die formale Definition verzichten. Exponentiation haben wir bereits für rationale Zahlen kennengelernt, wie diese Definition auf reelle Zahlen ausgeweitet werden kann überlassen wir der Mathematikvorlesung. Diese Funktion ist die Grundlage für Exponentialfunktionen zu einer beliebigen Basis  $a \in \mathbb{R}$ . Für eine graphische Veranschaulichung siehe Abbildung 4.11.

### 4.4.6. Logarithmusfunktionen

Die Logarithmusfunktion ist die Umkehrfunktion der Exponentialfunktion:

$$\begin{aligned}\log_a : \mathbb{R}^+ &\rightarrow \mathbb{R} \\ \log_a : x &\mapsto \log_a(x)\end{aligned}$$

Mit ihr kann man also die Operation des Potenzierens wieder rückgängig machen. Wieder verzichten wir auf die formale Definition.

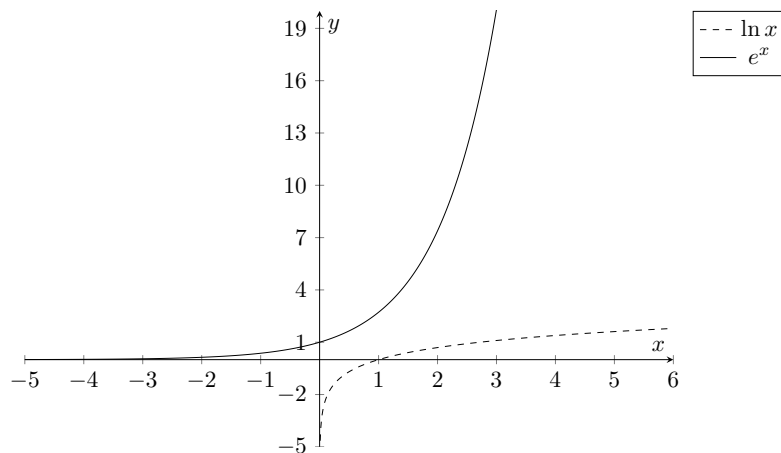


Abbildung 4.11.: Natürliche Exponentialfunktion und natürlicher Logarithmus.

## 4.5. Nullstellenberechnung

Das Finden von Nullstellen ist sehr wichtig für die Analysis, so gibt es z. B. Zusammenhänge zwischen den Nullstellen der Ableitungsfunktion  $f'(x)$  und der eigentlichen Funktion  $f$  (siehe Abschnitt 4.6).

**Definition 4.29** (Nullstelle). Sei  $f : \mathbb{R} \rightarrow \mathbb{R}$  eine Funktion. Als Nullstelle von  $f$  bezeichnet man eine Lösung der Gleichung  $f(x) = 0$ .

**Beispiel 4.30** (Nullstellen). Die Funktion  $f(x) = x^2 + x - 20$  hat die Nullstellen  $x_1 = -5, x_2 = 4$  (Abbildung 4.12a). Die Funktion  $g(x) = 5 \sin x$  hat die Nullstellen  $\pi \cdot a$  für  $a \in \mathbb{Z}$  (Abbildung 4.12b).

Im Folgenden wollen wir uns ausschließlich mit dem Finden von Nullstellen für *Polynome* beschäftigen. Aus der Schule sollte die *pq*-Formel bekannt sein (oder die allgemeinere Mitternachtsformel).

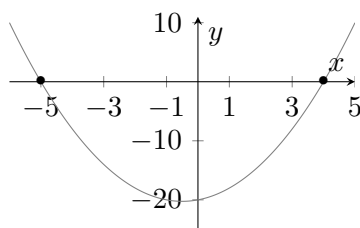
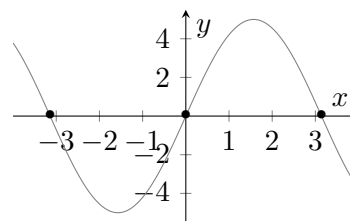
(a)  $f(x) = x^2 + x - 20$ (b)  $g(x) = 5 \sin x$ 

Abbildung 4.12.: Nullstellen von Funktionen.

#### 4. Analysis

**Satz 4.31** (*pq-Formel*). Sei  $f(x) = x^2 + px + q$  für  $p, q \in \mathbb{R}$ . Dann sind die Nullstellen von  $f$  gegeben durch

$$x_{1,2} = -\frac{p}{2} \pm \sqrt{\left(\frac{p}{2}\right)^2 - q}$$

**Hinweis.** Allgemeinere Gleichungen des Typs  $ax^2 + bx + c = 0$  kann man durch  $a$  teilen und so die *pq-Formel* verwenden.

Für Gleichungen dieser Form liefert auch die *Mitternachtsformel* die Nullstellen:

$$x_{1,2} = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a}$$

**Beispiel 4.32.** Sei  $f(x) = x^2 + 12x + 20$ . Dann sind die Nullstellen

$$\begin{aligned} x_{1,2} &= -\frac{12}{2} \pm \sqrt{\left(\frac{12}{2}\right)^2 - 20} \\ &= -6 \pm \sqrt{16} \\ x_1 &= -2 \\ x_2 &= -10 \end{aligned}$$

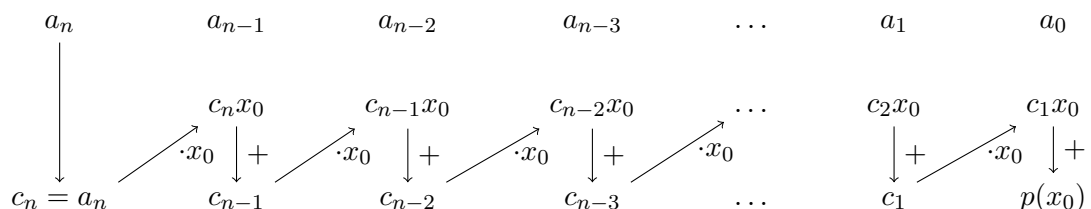
Diese Formel ist jedoch natürlich nur für Polynome vom Grad zwei sinnvoll. Für Polynome von höherem Grad ist es allgemein schwerer, Nullstellen zu finden. Es finden sich in einigen Formelsammlungen jedoch auch Formeln für Polynome von höherem Grad. Wir wollen an dieser Stelle das *Horner-Schema* vorstellen. Dieses kann verwendet werden, um ein Polynom in einer etwas anderen Darstellung zu repräsentieren, dabei bekommt man kleinere Polynome, deren Nullstellen einfacher zu bestimmen sind.

**Satz 4.33** (*Horner-Darstellung eines Polynoms*). Sei  $p : \mathbb{R} \rightarrow \mathbb{R}$  ein Polynom vom Grad  $n$ , also  $p(x) = \sum_{i=0}^n a_i \cdot x^i$ . Dann ist

$$\begin{aligned} p(x) &= a_0 + a_1x + \cdots + a_nx^n \\ &= a_0 + x(a_1 + x(a_2 + \cdots + x(a_{n-1} + x \cdot a_n) \dots)) \end{aligned}$$

die *Horner-Darstellung* des Polynoms.

**Satz 4.34** (*Horner-Schema*). Gegeben sei ein Polynom  $p(x) = \sum_{i=0}^n a_i \cdot x^i$  und  $x_0 \in \mathbb{R}$ . Dann lässt sich  $p(x_0)$  mit dem Horner-Schema wie folgt berechnen:



Zur Berechnung geht man also folgendermaßen vor:

Man füllt die erste Zeile mit den Koeffizienten des Polynoms. In die erste Spalte der untersten Zeile schreibt man den Koeffizienten  $a_n$ . Dann multipliziert man den Eintrag mit  $x_0$  und schreibt ihn in die mittlere Zeile der zweiten Spalte. Daraufhin addiert man diesen Eintrag mit dem Koeffizienten der Spalte und schreibt die Summe in die unterste Zeile. Schließlich wiederholt man diesen Vorgang und erhält in der untersten Zeile der letzten Spalte  $p(x_0)$ .

Nach der Berechnung gilt:

$$p(x) = (x - x_0) \sum_{i=1}^n c_i \cdot x^{i-1} + p(x_0)$$

Zur Auswertung des Polynoms werden weniger Multiplikationen als bei trivialem Einsetzen benötigt.

**Beispiel 4.35** (Anwendung des Horner-Schemas). Sei  $p(x) = x^4 - 8x^3 - 3x^2 + 62x + 56$ . Wir wollen das Polynom an der Stelle  $x_0 = 2$  auswerten. Die einzelnen Schritte sind in folgender Tabelle dargestellt:

|   |    |     |     |     |
|---|----|-----|-----|-----|
| 1 | -8 | -3  | 62  | 56  |
|   | 2  | -12 | -30 | 64  |
| 1 | -6 | -15 | 32  | 120 |

Wir erhalten, dass  $p(2) = 120$  und  $p(x) = (x - 2)(x^3 - 6x^2 - 15x + 32) + 120$ .

Besonders interessant ist der Fall, dass  $x_0$  eine Nullstelle des Polynoms ist. In diesem Fall erhalten wir  $p(x) = (x - x_0) \cdot \hat{p}(x)$ . Es gilt  $p(x) = 0$  gdw.  $x = x_0$  oder  $\hat{p}(x) = 0$ , wobei  $\hat{p}(x)$  ein Polynom kleineren Grades ist. Das Hornerverfahren entspricht in diesem Fall also der Polynomdivision  $p(x) : (x - x_0) = \hat{p}(x)$ .

**Hinweis.** Um das Horner-Schema zur Berechnung von Nullstellen anwenden zu können, muss man erst eine Nullstelle „raten“. Möchte man rationale Nullstellen eines Polynoms raten, so hilft der folgende Satz: Für jede rationale Nullstelle  $x' = \frac{u}{v}$  eines ganzzahligen Polynoms  $p$  gilt, dass  $a_0$  durch  $u$  und  $a_n$  durch  $v$  teilbar ist.

## 4.6. Differentialrechnung

Denen die fragen, was die unendlich kleine Größe in der Mathematik wäre, antworten wir, dass sie eigentlich Null sei.

---

*Leonhard Euler*

Ein wesentlicher Bereich der Analysis für Ingenieur\*innen und Informatiker\*innen ist die Differentialrechnung. Sie beschäftigt sich mit der Beschreibung lokaler Veränderungen

## 4. Analysis

von Funktionen. Der grundlegende Begriff der *Ableitung* beschreibt diese lokale Änderung. Ableitungen spielen in vielen Bereichen eine Rolle, in denen man natürliche Phänomene beschreibt. So ist etwa in der Physik die Veränderung der Geschwindigkeit, alias Beschleunigung, mathematisch betrachtet die Ableitung der Geschwindigkeit (nach der Zeit).

### Ableitung

Als Ableitung einer Funktion in einem bestimmten Punkt versteht man die Steigung der Tangente<sup>5</sup> in diesem Punkt. Bei linearen Funktionen ist so eine Tangente natürlich schnell gefunden. Dort hatten wir zur Bestimmung der Steigung zwei Punkte der Funktion benutzt und den Quotienten  $\frac{\Delta y}{\Delta x}$  gebildet. Dieser ist für lineare Funktionen an allen Stellen gleich und schon haben wir unsere Steigung der Tangente. Bereits für Polynome vom Grad höher als 1 wird es jedoch schwierig, eine solche Steigung zu finden. Natürlich können wir nicht einfach für die linearen Funktionen eine Ausnahme bilden und die Steigung für alle anderen Funktionen komplett anders definieren. Um daher weiterhin von unseren zwei Punkten Gebrauch zu machen, müssen wir zunächst den Begriff des Grenzwertes ein wenig anpassen.

Bisher haben wir den Grenzwert bei Folgen genutzt, um ihren Wert „im Unendlichen“ zu bestimmen. Eine allgemeinere Definition des Grenzwertbegriffs lässt sich jedoch auch auf Funktionen  $f : \mathbb{R} \rightarrow \mathbb{R}$  sinnvoll anwenden. Die formale Definition wollen wir euch an dieser Stelle jedoch ersparen, sie wird in der Mathematik I aber sicher eingeführt. Stattdessen wollen wir bei einer intuitiven Beschreibung bleiben:

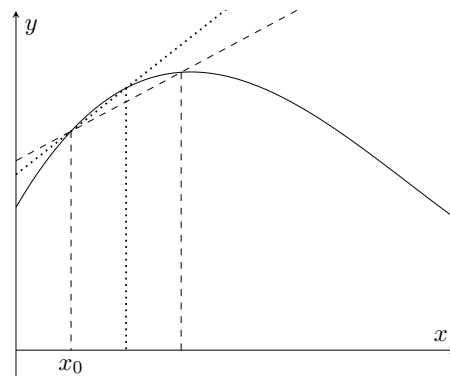
**Definition 4.36** (Grenzwert einer Funktion). Den Grenzwert einer Funktion  $f$  bezüglich einer Stelle  $x_0$  erhält man, wenn man in  $f$  Werte einsetzt, die sich immer mehr an  $x_0$  annähern,  $x_0$  jedoch nie ganz erreichen. Man schreibt dann  $\lim_{x \rightarrow x_0} f(x)$ .

Dabei muss man jedoch beachten, dass es einen Unterschied machen kann, ob man sich „von rechts“ oder „von links“ an dieses  $x_0$  annähert. Betrachtet man beispielsweise die Funktion  $f(x) = \frac{1}{x}$ , so wird man feststellen, dass  $\lim_{x \rightarrow 0} \frac{1}{x} = \infty$  ist, wenn man sich „von rechts“ annähert, jedoch  $-\infty$ , wenn man sich „von links“ annähert.

Nun kann man die Steigung in einem Punkt  $x_0$  bestimmen, indem man als zweiten Punkt für die Bestimmung der Steigung einen Punkt  $x_0 + \varepsilon$  wählt, der ein kleines Stückchen weit von  $x_0$  weg ist (siehe Abbildung 4.13 und 4.14). Die Steigung dieser Geraden ist natürlich nicht exakt die gleiche wie die von der Tangente  $f(x_0)$ . Jedoch nähert sich unsere Gerade immer mehr an die tatsächliche Tangente an, je kleiner wir  $\varepsilon$  wählen. Lassen wir nun  $\varepsilon$  gegen 0 streben, so erhalten wir die tatsächliche Steigung der Tangente in  $f(x_0)$ .

Da für die Berechnung der Steigung in der Stelle  $x_0$  nun der Quotient aus zwei (infinitesimal kleinen) Differenzen gebildet werden muss, definiert man den Differentialquotient.

<sup>5</sup>Nicht zu verwechseln mit einem veralgten Seevogel.



**Abbildung 4.13.:** Ein Polynom dritten Grades und zwei Geraden, die die Steigung an der Stelle  $x_0$  annähern. Je kleiner der Abstand zwischen den zwei Punkten ist, desto genauer entspricht die Gerade der tatsächlichen Tangente.

**Definition 4.37** (Differentialquotient). Sei  $f$  eine Funktion, die auf einem Intervall  $I \subseteq \mathbb{R}$  definiert ist und sei  $x_0 \in I$ . Dann nennt man die Funktion  $f$  in  $x_0$  *differenzierbar*, falls der Grenzwert

$$\frac{df(x_0)}{dx} := \lim_{\varepsilon \rightarrow 0} \frac{f(x_0 + \varepsilon) - f(x_0)}{\varepsilon}$$

existiert. Ist  $f$  in jedem Punkt  $x_0 \in I$  differenzierbar, so sagt man, dass die Funktion auf  $I$  differenzierbar ist. Die Schreibweise mit  $d$  geht auf *Carl Friedrich Gauß*<sup>6</sup> zurück, man spricht „d  $f$  von  $x_0$  nach d  $x$ “.

**Hinweis.** Eine Funktion, deren Ableitung stetig ist, nennt man auch *stetig differenzierbar*, nicht zu verwechseln mit *stetig und differenzierbar*, was über die Stetigkeit der Ableitung keine Aussage trifft.

**Definition 4.38** (Ableitung). Sei  $f$  auf  $I \subseteq \mathbb{R}$  differenzierbar. Die Funktion  $f' : I \rightarrow \mathbb{R}$  mit

$$f'(x) := \frac{df(x)}{dx}$$

heißt *Ableitung* von  $f$ .

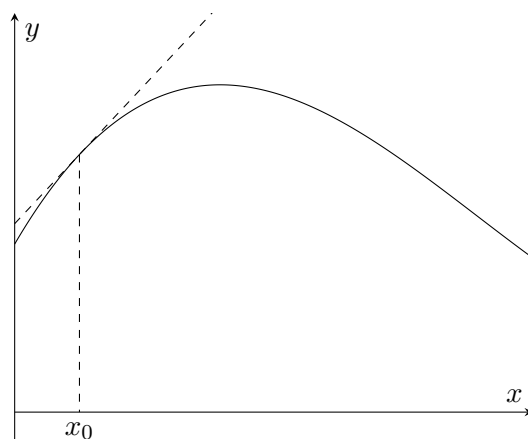
#### 4.6.1. Gängige Ableitungen

- *Polynome* werden differenziert, indem jeder Koeffizient  $a_i$  mit  $i$  multipliziert und die korrespondierende Potenz um eins verringert wird. Formal schreibt sich das wie folgt:

$$\text{Für } f(x) = \sum_{i=0}^n a_i \cdot x^i \text{ ist die Ableitung } f'(x) = \sum_{i=0}^n i \cdot a_i \cdot x^{i-1}$$

<sup>6</sup>deutscher Mathematiker, 1777 - 1855.

#### 4. Analysis



**Abbildung 4.14.:** Tangente, die in der Stelle  $x_0$  anliegt.

| $f(x)$    | $f'(x)$       |
|-----------|---------------|
| $\sin(x)$ | $\cos(x)$     |
| $\cos(x)$ | $-\sin(x)$    |
| $\ln(x)$  | $\frac{1}{x}$ |
| $e^x$     | $e^x$         |

**Tabelle 4.2.:** Weitere gängige Ableitungen.

**Beispiel 4.39.** Sei  $f(x) = \frac{3}{2}x^2 + 2x + 1$ , dann ist die Ableitung  $f'(x) = 2 \cdot \frac{3}{2}x^{2-1} + 1 \cdot 2x^{1-1} + 0 \cdot 1x^{0-1} = 3x + 2$

- Die Ableitung der *Wurzelfunktion* funktioniert entsprechend der Ableitung der Polynome:

Für  $f(x) = \sqrt{x} = x^{\frac{1}{2}}$  ist die Ableitung  $f'(x) = \frac{1}{2}x^{-\frac{1}{2}} = \frac{1}{2\sqrt{x}}$

- Tabelle 4.2 enthält weitere wichtige Ableitungen.

Die Herleitung der Ableitungsregeln überlassen wir der Mathe 1 Vorlesung und wollen nur ein Beispiel betrachten.

**Beispiel 4.40** (Herleitung einer Ableitungsfunktion). Den Regeln folgend ist die Ableitungsfunktion von  $f(x) = x^2$  gegeben durch  $f'(x) = 2x$ .

$$\begin{aligned}
 f'(x) &= \lim_{\varepsilon \rightarrow 0} \frac{(x + \varepsilon)^2 - x^2}{\varepsilon} \\
 &= \lim_{\varepsilon \rightarrow 0} \frac{x^2 + \varepsilon^2 + 2x\varepsilon - x^2}{\varepsilon} \\
 &= \lim_{\varepsilon \rightarrow 0} \frac{\varepsilon^2 + 2x\varepsilon}{\varepsilon} \\
 &= \lim_{\varepsilon \rightarrow 0} \frac{\varepsilon^2}{\varepsilon} + \frac{2x\varepsilon}{\varepsilon} \\
 &= \lim_{\varepsilon \rightarrow 0} \varepsilon + 2x \\
 &= 2x
 \end{aligned}$$

#### 4.6.2. Rechenregeln für die Differentiation

Da man nicht immer nur eine reine Funktion aus einer Funktionsklasse ableiten muss, sondern auch Kombinationen davon, benötigt man weitere Ableitungsregeln, die hier vorgestellt werden.

**Satz 4.41** (Faktorenregel). Sei  $f$  eine Funktion und  $c \in \mathbb{R}$  eine Konstante. Dann ist

$$(c \cdot f(x))' = c \cdot f'(x)$$

*Beweis.*

$$\begin{aligned}
 (c \cdot f(x))' &= \lim_{\varepsilon \rightarrow 0} \frac{c \cdot f(x + \varepsilon) - c \cdot f(x)}{\varepsilon} \\
 &= \lim_{\varepsilon \rightarrow 0} \frac{c \cdot (f(x + \varepsilon) - f(x))}{\varepsilon} \\
 &= c \cdot \lim_{\varepsilon \rightarrow 0} \left( \frac{f(x + \varepsilon) - f(x)}{\varepsilon} \right) \\
 &= c \cdot f'(x)
 \end{aligned}$$

□

**Beispiel 4.42.** Sei  $f(x) = 5 \cdot \cos x$ . Dann ist

$$\begin{aligned}
 f'(x) &= 5 \cdot (\cos x)' \\
 &= -5 \cdot \sin x
 \end{aligned}$$

**Satz 4.43** (Linearität der Ableitung). Die Ableitung ist linear, das bedeutet, wenn  $f(x) = g(x) + h(x)$  die Addition zweier Funktionen ist, dann ist die Ableitung

$$f'(x) = g'(x) + h'(x)$$

#### 4. Analysis

**Satz 4.44** (Produktregel). Sei  $f(x) = g(x) \cdot h(x)$  das Produkt zweier Funktionen, dann ist die Ableitung  $f'(x) = g(x) \cdot h'(x) + g'(x) \cdot h(x)$ .

**Beispiel 4.45.**

$$\begin{aligned}f(x) &= \ln(x) \cdot x \\g(x) &= \ln(x) \\g'(x) &= \frac{1}{x} \\h(x) &= x \\h'(x) &= 1 \\f'(x) &= \ln(x) \cdot 1 + \frac{1}{x} \cdot x \\&= \ln(x) + 1\end{aligned}$$

**Satz 4.46** (Kettenregel). Sei  $f$  die Verkettung zweier Funktionen  $g, h$ , also  $f(x) = g(h(x))$ . Wenn  $g, h$  differenzierbar sind, dann auch  $f(x)$  mit

$$f'(x) = g'(h(x)) \cdot h'(x)$$

**Beispiel 4.47.** Sei  $f(x) = \cos(x^2 + \ln x)$ . Dann ist  $f(x) = g(h(x))$  mit  $g(y) = \cos y$  und  $h(x) = x^2 + \ln x$ .

Dann ist

$$\begin{aligned}g'(y) &= -\sin y \\h'(x) &= 2x + \frac{1}{x}\end{aligned}$$

also

$$f'(x) = -\sin(x^2 + \ln x) \cdot \left(2x + \frac{1}{x}\right)$$

für  $x > 0$ .

**Beispiel 4.48.** Sei  $f(x) = e^{2x} + \cos(\sin(x^2))$ .

Dann müssen wir zuerst  $e^{2x}$  ableiten. Hierbei handelt es sich um eine verkettete Funktion mit  $g_1(y) = e^y$  und  $h_1(x) = 2x$ . Damit erhalten wir:

$$\begin{aligned}g_1'(y) &= e^y \\h_1'(x) &= 2 \\(e^{2x})' &= 2e^{2x}\end{aligned}$$

Dann müssen wir  $\cos(\sin(x^2))$  ableiten. Dies ist eine verkettete Funktion mit

$$\begin{aligned}g_2(z) &= \cos z \\h_2(y) &= \sin y^2 \\g_2'(z) &= -\sin z\end{aligned}$$

Wir benötigen also die Ableitung von  $\sin x^2$ . Dies ist wieder eine verkettete Funktion mit

$$\begin{aligned}g_3(y) &= \sin y \\h_3(x) &= x^2 \\g'_3(y) &= \cos y \\h'_3(x) &= 2x \\(\sin(x^2))' &= \cos(x^2) \cdot 2x\end{aligned}$$

und damit

$$\left(\cos\left(\sin\left(x^2\right)\right)\right)' = -\sin\left(\sin\left(x^2\right)\right) \cdot \cos\left(x^2\right) \cdot 2x$$

Also ist

$$f'(x) = 2e^{2x} - \sin\left(\sin\left(x^2\right)\right) \cdot \cos\left(x^2\right) \cdot 2x$$

**Satz 4.49** (Quotientenregel). Sei  $f$  der *Quotient* zweier Funktionen  $g, h$ , also  $f(x) = \frac{g(x)}{h(x)}$ . Wenn  $g, h$  differenzierbar sind, dann auch  $f(x)$  mit

$$f'(x) = \frac{g'(x) \cdot h(x) - g(x) \cdot h'(x)}{(h(x))^2}$$

**Beispiel 4.50.** Sei  $f(x) = \frac{1}{x}$ , dann ist offenbar  $g(x) = 1$  und  $h(x) = x$ . Dann ist

$$\begin{aligned}g'(x) &= 0 \\h'(x) &= 1\end{aligned}$$

und damit

$$\begin{aligned}f'(x) &= \frac{g'(x) \cdot h(x) - g(x) \cdot h'(x)}{(h(x))^2} \\&= \frac{0 \cdot x - 1 \cdot 1}{x^2} \\&= \frac{-1}{x^2}\end{aligned}$$

#### 4. Analysis

**Beispiel 4.51.** Sei  $f(x) = \tan(x)$ . Diese Funktion können wir umformen zu  $f(x) = \frac{\sin(x)}{\cos(x)}$ . Nun können wir die Quotientenregel anwenden:

$$\begin{aligned} f'(x) &= \left( \frac{\sin(x)}{\cos(x)} \right)' \\ &= \frac{\cos(x) \cdot \cos(x) - \sin(x) \cdot (-\sin(x))}{(\cos(x))^2} \\ &= \frac{\cos^2(x) + \sin^2(x)}{\cos^2(x)} \\ &= \frac{\cos^2(x)}{\cos^2(x)} + \frac{\sin^2(x)}{\cos^2(x)} \\ &= 1 + \tan^2(x) \end{aligned}$$

Alternativ kann man bei der Umformung benutzen, dass  $\cos^2(x) + \sin^2(x) = 1$  ist:

$$\begin{aligned} f'(x) &= \frac{\cos^2(x) + \sin^2(x)}{\cos^2(x)} \\ &= \frac{1}{\cos^2(x)} \end{aligned}$$

### 4.7. Integralrechnung

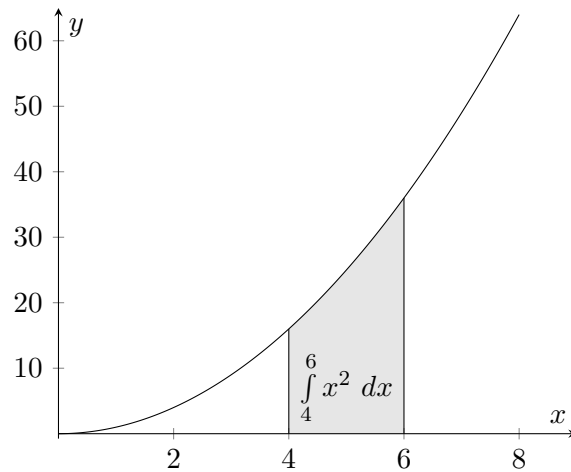
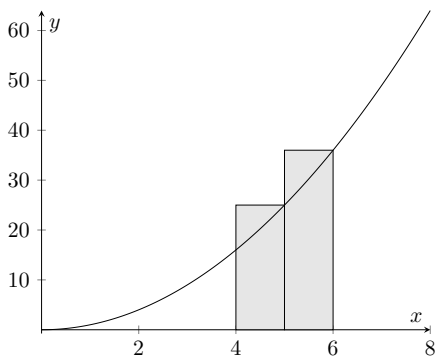
Der andere große Bereich der Analysis ist die Integralrechnung. Wir werden sehen, dass zwischen Integral- und Differentialrechnung ein enger Zusammenhang besteht. Die Integralrechnung befasst sich mit der Berechnung von Flächen zwischen  $x$ -Achse und Funktionsgraphen.

**Beispiel 4.52** (Motivation). Wir wollen die Fläche zwischen der  $x$ -Achse und dem Graph der Funktion  $f(x) = x^2$  zwischen den Stellen  $x_0 = 4$  und  $x_1 = 6$  berechnen (siehe Abbildung 4.15).

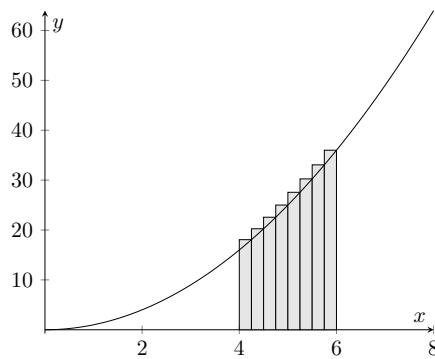
Die Idee, um Flächen unter Funktionsgraphen zu berechnen, ist es, die entsprechende Fläche durch *Rechtecke* anzunähern. In Abbildung 4.16a ist die Fläche unter dem Funktionsgraphen angenähert durch zwei Rechtecke. Beide Rechtecke haben eine Breite von 1. Die Höhe der Rechtecke ergibt sich aus den Funktionswerten  $f(5)$  und  $f(6)$ .

Ähnlich wie bei der Differentiation mit einer Geraden durch zwei Punkte ist dies jedoch nur eine verhältnismäßig grobe Näherung. Und ganz genauso wie wir bei der Differentiation den Abstand der Punkte immer weiter verringert haben, können wir hier die Genauigkeit der Näherung erhöhen, indem wir immer mehr Rechtecke einfügen (siehe Abbildung 4.16b).

**Definition 4.53** (Bestimmtes Integral). Sei  $f$  eine auf dem Intervall  $[a; b]$  definierte Funktion, die auf diesem Intervall beschränkt ist. Dann heißt  $\int_a^b f(x) dx$  das *bestimmte Integral* über  $[a; b]$  und ist wie folgt definiert.

Abbildung 4.15.: Fläche unter  $x^2$  zwischen den Stellen 4 und 6.

(a) Annäherung durch zwei Rechtecke.



(b) Annäherung durch acht Rechtecke.

Abbildung 4.16.: Annäherung der Fläche unter der Funktion  $f(x) = x^2$  durch Rechtecke zwischen den Stellen 4 und 6.

Sei  $a = x_0 < x_1 < \dots < x_n = b$  eine *Einteilung* des Intervalls  $[a; b]$ .

$$\int_a^b f(x) \, dx = \lim_{n \rightarrow \infty} \sum_{i=1}^n f(\zeta_i) \cdot (x_i - x_{i-1})$$

wobei  $x_{i-1} \leq \zeta_i \leq x_i$ .

Ist  $b < a$  so definiert man  $\int_a^b f(x) \, dx := -\int_b^a f(x) \, dx$ .

**Hinweis.** Unter dem Intervall  $[a; b]$  verstehen wir die folgende Menge reeller Zahlen:  
 $[a; b] := \{x \in \mathbb{R} \mid a \leq x \leq b\}$ .

**Hinweis.** Das Integral kann durchaus auch einen *negativen* Wert annehmen, zum Beispiel wenn die Funktion unterhalb der  $x$ -Achse verläuft. Man kann also nicht direkt die Fläche

#### 4. Analysis

| $f(x)$        | $F(x)$                 |
|---------------|------------------------|
| $x^n$         | $\frac{1}{n+1}x^{n+1}$ |
| $\sin(x)$     | $-\cos(x)$             |
| $\cos(x)$     | $\sin(x)$              |
| $\frac{1}{x}$ | $\ln(x)$               |
| $e^x$         | $e^x$                  |

**Tabelle 4.3.:** Einige wichtige Stammfunktionen.

zwischen Kurve und  $x$ -Achse ablesen, vielmehr muss man vorher alle Schnittpunkte der Funktion mit der  $x$ -Achse bestimmen und die entsprechenden Teilflächen aufsummieren.

**Definition 4.54** (Stammfunktion). Eine Funktion  $F$  ist Stammfunktion von  $f$ , falls  $F' = f$ .

**Definition 4.55** (Unbestimmtes Integral). Das *unbestimmte Integral*

$$\int f(x) \, dx := \{F(x) \mid F' = f\}$$

ist die Menge aller Stammfunktionen einer Funktion  $f$ . Alle diese Stammfunktionen unterscheiden sich lediglich durch einen konstanten Summanden  $C$ .

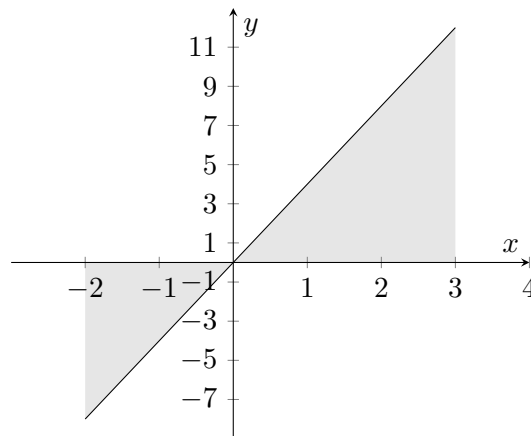
Es gilt folgender Zusammenhang: Ist  $F$  eine beliebige Stammfunktion von  $f$ , so ist

$$\int_a^b f(x) \, dx = F(b) - F(a)$$

Man schreibt für  $F(b) - F(a)$  oft abkürzend auch  $F(x) \Big|_a^b$  oder  $[F(x)]_a^b$ .

Im Folgenden verwenden wir das unbestimmte Integral um eine beliebige Stammfunktion (üblicherweise mit  $C = 0$ ) zu bezeichnen.

**Beispiel 4.56** (Flächenbestimmung). Sei  $f(x) = 4x$ . Wir wollen die Fläche zwischen  $a = -2$  und  $b = 3$  berechnen. Zuerst stellen wir fest, dass die Funktion einen Vorzeichenwechsel bei  $x_0 = 0$  hat: vorher sind die Funktionswerte negativ und ab  $x_0$  positiv. Zudem ist 0 die einzige Nullstelle der Funktion, vgl. auch Abbildung 4.17.



**Abbildung 4.17.:** Um die Fläche der Funktion  $f(x) = 4x$  zwischen  $a = -2$  und  $b = 3$  berechnen zu können, müssen wir den Vorzeichenwechsel berücksichtigen.

Wir berechnen:

$$\begin{aligned} \int 4x \, dx &= 2x^2 \\ \int_{-2}^0 4x \, dx &= 2 \cdot 0^2 - 2 \cdot (-2)^2 = -8 \\ \int_0^3 4x \, dx &= 2 \cdot 3^2 - 2 \cdot 0^2 = 18 \end{aligned}$$

Für die Gesamtfläche  $A$  addieren wir die Beträge der entsprechenden Werte und erhalten  $A = |-8| + |18| = 26$ .

#### 4.7.1. Integrationsregeln

**Satz 4.57** (Faktorregel). Sei  $f$  eine Funktion und  $c \in \mathbb{R}$  eine Konstante. Dann ist

$$\int c \cdot f(x) \, dx = c \int f(x) \, dx$$

**Satz 4.58** (Linearität der Integration). Die Integration ist linear, das bedeutet, wenn  $f(x) = g(x) + h(x)$  die Summe zweier Funktionen ist, und  $G$  und  $H$  die Stammfunktionen von  $g$  und  $h$  sind, dann ist deren Stammfunktion

$$F(x) = G(x) + H(x)$$

**Satz 4.59** (Partielle Integration). Sind  $f, g$  differenzierbare Funktionen, so gilt

$$\int f'(x) \cdot g(x) \, dx = f(x) \cdot g(x) - \int f(x) \cdot g'(x) \, dx$$

#### 4. Analysis

Um die partielle Integration anwenden zu können, muss also eine der Funktionen abgeleitet werden und die andere Funktion integriert. Man versucht meist eine Funktion zu finden, die sich leicht integrieren lässt und hofft, dass das resultierende Integral einfacher zu lösen ist.

**Beispiel 4.60** (Partielle Integration).

$$\int x \ln x \, dx$$

Wähle nun  $f'(x) = x$  und  $g(x) = \ln x$ . Integrieren bzw. Ableiten liefert  $f(x) = \frac{x^2}{2}$  und  $g'(x) = \frac{1}{x}$ . Damit ergibt sich:

$$\begin{aligned} \int x \ln x \, dx &= \frac{x^2 \ln x}{2} - \int \frac{x^2}{2} \cdot \frac{1}{x} \, dx \\ &= \frac{x^2 \ln x}{2} - \int \frac{x}{2} \, dx \\ &= \frac{x^2 \ln x}{2} - \frac{1}{2} \int x \, dx \\ &= \frac{x^2 \ln x}{2} - \frac{1}{2} \cdot \frac{x^2}{2} \\ &= \frac{x^2 \ln x}{2} - \frac{x^2}{4} \end{aligned}$$

**Beispiel 4.61.**

$$\int \ln x \, dx$$

Wähle nun  $f'(x) = 1$  und  $g(x) = \ln x$  und damit  $f(x) = x$  und  $g'(x) = \frac{1}{x}$ . Dann ergibt sich

$$\begin{aligned} \int \ln x \, dx &= x \ln x - \int x \cdot \frac{1}{x} \, dx \\ &= x \ln x - x \\ &= x (\ln x - 1) \end{aligned}$$

**Beispiel 4.62.** In diesem Beispiel werden wir uns eines Tricks bedienen, den man sich zum Integrieren merken sollte!

$$\int \sin x \cdot \cos x \, dx$$

Wähle  $f'(x) = \cos x$  und  $g(x) = \sin x$ . Dann ist  $f(x) = \sin x$  und  $g'(x) = \cos x$ . Dann erhalten wir

$$\int \sin x \cdot \cos x \, dx = \sin^2 x - \int \sin x \cdot \cos x \, dx$$

Nun müssen wir allerdings wieder das Integral  $\int \sin x \cos x \, dx$  lösen, was hat uns der Schritt also gebracht? Eine einfache Umformung bringt uns weiter:

$$2 \int \sin x \cdot \cos x \, dx = \sin^2 x$$

also folgt mit Division durch 2

$$\int \sin x \cdot \cos x \, dx = \frac{\sin^2 x}{2}$$

Ähnlich zur Kettenregel bei der Differentiation gibt es auch für die Integration die Möglichkeit, verkettete Funktionen mit Hilfe bereits bekannter Stammfunktionen zu berechnen. Dies erfolgt durch Substitution, ist allerdings, wie die partielle Integration auch, etwas komplizierter.

**Satz 4.63** (Integration durch Substitution). Sind  $f, \varphi$  Funktionen und  $\varphi$  stetig differenzierbar, sowie  $t = \varphi(x)$ , so gilt

$$\int_a^b f(\varphi(x)) \cdot \varphi'(x) \, dx = \int_{\varphi(a)}^{\varphi(b)} f(t) \, dt$$

Um die Substitutionsregel anwenden zu können, muss also die Ableitung der inneren Funktion ebenfalls im Integral stehen. Zu beachten ist, dass auch die Grenzen des Integrals substituiert werden müssen.

Für das unbestimmte Integral gilt:

$$\int f(x) \, dx = \int \underbrace{f(g(u))}_{f(x)} \cdot \underbrace{g'(u)}_{dx} \, du$$

wenn  $g$  eine stetig differenzierbare Funktion ist und für ihre Ableitung  $g'$  überall ungleich 0 ist. Die substituierte Funktion  $g(u)$  muss dabei umkehrbar sein.

Der Trick besteht darin, den Ausdruck der Form  $f(x) \, dx$  in die „passende Form“ zu bringen.

**Beispiel 4.64** (Integration durch Substitution).

$$\begin{aligned} \int_{\frac{3}{2}}^2 \cos(\pi x) \, dx &= \int_{\frac{3}{2}}^2 \cos(\pi x) \, dx \\ \text{(Substitutionsregel)} &= \int_{\frac{3\pi}{2}}^{2\pi} \frac{1}{\pi} \cos(t) \, dt \\ \text{(Faktorregel)} &= \frac{1}{\pi} \cdot \int_{\frac{3\pi}{2}}^{2\pi} \cos(t) \, dt \\ \text{(Integration)} &= \frac{1}{\pi} [\sin(t)]_{\frac{3\pi}{2}}^{2\pi} \\ &= \frac{1}{\pi} \left( \sin(2\pi) - \sin\left(\frac{3\pi}{2}\right) \right) \\ &= \frac{1}{\pi} (0 + 1) = \frac{1}{\pi} \end{aligned}$$

**Beispiel 4.65.**

$$\begin{aligned}\int \cos 4x \, dx &= \frac{1}{4} \cdot \int \cos 4x \cdot 4 \, dx \\ &= \frac{1}{4} \cdot \int \cos u \, du \\ &= \frac{1}{4} \cdot \sin u \\ &= \frac{1}{4} \sin 4x\end{aligned}$$

In diesem Beispiel kann man sehen, wie die Funktion auf die richtige Form gebracht wird. Die eigentliche Substitution besteht aus  $x = g(u) = \frac{u}{4}$ . Da im hinteren Teil die Ableitung von  $g$ , also  $g'(u) = \frac{1}{4}$  benötigt wird, erweitern wir mit  $\frac{1}{4}$  und erhalten die gewünschte Form.

**Beispiel 4.66.** Ein wichtiger Spezialfall der Substitutionsregel ist das folgende Integral:

$$\int \frac{f'(x)}{f(x)} \, dx = \ln |f(x)|$$

Das Integral einer Ableitung einer Funktion geteilt durch die Funktion selbst ist also *immer* ihr natürlicher Logarithmus.

## 5. Lineare Algebra

Menschen, die von der Algebra nichts wissen, können sich auch nicht die wunderbaren Dinge vorstellen, zu denen man mit Hilfe der genannten Wissenschaft gelangen kann.

---

*Gottfried Wilhelm Leibniz*

Ein weiteres großes Gebiet der Mathematik ist die lineare Algebra oder Vektoralgebra. Da die analytische Geometrie ein großer Anwendungsbereich mit sehr schönen Beispielen ist, werden wir viele Veranschaulichungen aus der Geometrie verwenden.

### 5.1. Vektoren

In den vergangenen Kapiteln haben wir uns hauptsächlich mit „einfachen“ Zahlenräumen wie  $\mathbb{N}$ ,  $\mathbb{R}$  usw. auseinandergesetzt. In der linearen Algebra beschäftigen wir uns nun aber mit Tupeln. Wir werden uns hier mit den sogenannten Vektorräumen  $\mathbb{R}^n$  befassen.

#### 5.1.1. Nomenklatur

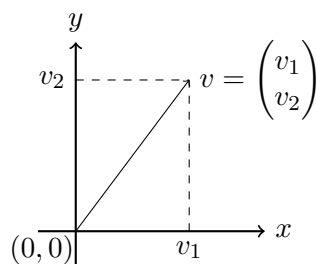
**Definition 5.1** (Vektor). Unter einem *Vektor*  $v$  versteht man ein  $n$ -Tupel von Zahlen. Dabei ist  $n$  die *Dimension* eines Vektors. Das  $i$ -te Element des Vektors wird mit  $v_i$  bezeichnet. Man unterscheidet je nach Schreibweise zwischen Spalten- und Zeilenvektoren, die sich durch Transposition ineinander umwandeln lassen. Man schreibt:

$$v = \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{pmatrix} = \begin{pmatrix} v_1 & v_2 & \dots & v_n \end{pmatrix}^T$$

Geometrisch versteht man unter dem Begriff ein Objekt, welches eine Parallelverschiebung in der Ebene oder im Raum beschreibt. Anschaulich lässt sich ein Vektor dann als ein Pfeil auffassen, der den Nullpunkt mit dem Bildpunkt verbindet. Siehe Abbildung 5.1.

Im Sinne der Übersichtlichkeit einigen wir uns an dieser Stelle auf die Bezeichnungen von Variablen wie in Tabelle 5.1 angegeben.

## 5. Lineare Algebra



**Abbildung 5.1.:** Ein Vektor im Vektorraum  $\mathbb{R}^2$ .

| Typ              | Variablenname               | Beispiel                                           |
|------------------|-----------------------------|----------------------------------------------------|
| Skalar<br>(Zahl) | Griechische Kleinbuchstaben | $\lambda = 3$                                      |
| Vektor           | Römische Kleinbuchstaben    | $v = \begin{pmatrix} 2 \\ 3 \end{pmatrix}$         |
| Matrix           | Römische Großbuchstaben     | $A = \begin{pmatrix} 2 & 3 \\ 1 & 0 \end{pmatrix}$ |

**Tabelle 5.1.:** Unsere Namenskonvention für Variablen.

### 5.1.2. Vektoroperationen

Alle nachfolgenden Operationen können sinngemäß auch auf Zeilenvektoren angewandt werden. Da jedoch Spaltenvektoren geläufiger sind, beschreiben wir die Operationen hier nur für Spaltenvektoren.

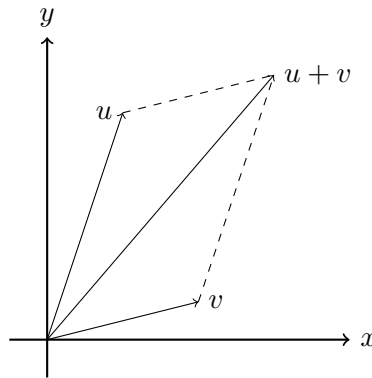
**Definition 5.2** (Betrag eines Vektors). Der *Betrag* oder auch die (*euklidische*) *Norm* eines Vektors  $\|v\|$  ist definiert als:

$$\left\| \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{pmatrix} \right\| := \left| \sqrt{v_1^2 + v_2^2 + \cdots + v_n^2} \right| = \left| \sqrt{\sum_{i=1}^n v_i^2} \right|$$

Der Betrag lässt sich einfach aus dem Satz des Pythagoras<sup>1</sup> herleiten und beschreibt die geometrische Länge eines Vektors.

**Hinweis.** Als *normierten Vektor* bezeichnet man einen Vektor mit der Länge 1.

<sup>1</sup>toter griechischer Mathematiker, ca. 570 - 495 v.Chr.

Abbildung 5.2.: Addition zweier Vektoren  $u, v$  im  $\mathbb{R}^2$ .

**Definition 5.3** (Addition zweier Vektoren). Seien  $u, v$  zwei Vektoren gleicher Dimension  $n$ . Dann werden diese addiert, indem ihre Komponenten miteinander addiert werden:

$$u + v = \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_n \end{pmatrix} + \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{pmatrix} := \begin{pmatrix} u_1 + v_1 \\ u_2 + v_2 \\ \vdots \\ u_n + v_n \end{pmatrix}$$

Die Addition von Vektoren ist *kommutativ* und *assoziativ*. Wir können diese als ein „Aneinandersetzen“ der Pfeile verstehen, siehe Abbildung 5.2.

**Definition 5.4** (Multiplikation mit einem Skalar). Sei  $\lambda$  ein Skalar und  $v$  ein Vektor, dann werden diese multipliziert, indem das Skalar mit jeder Komponente des Vektors multipliziert wird.

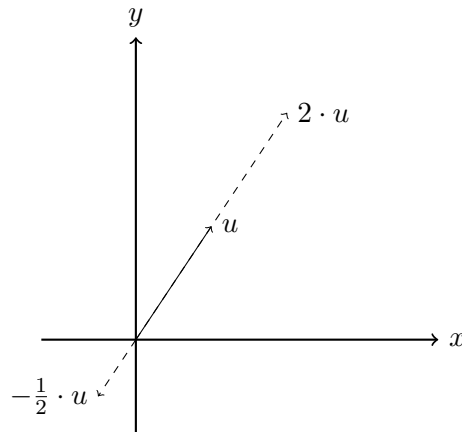
$$\lambda v = \lambda \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{pmatrix} := \begin{pmatrix} \lambda \cdot v_1 \\ \lambda \cdot v_2 \\ \vdots \\ \lambda \cdot v_n \end{pmatrix}$$

Die Multiplikation mit einem Skalar ist *kommutativ* und *assoziativ*.

**Hinweis.** Durch die Multiplikation mit einem Skalar wird die Länge des Vektors skaliert. Der Vektor zeigt in dieselbe Richtung, falls das Skalar positiv ist, sonst in die Gegenrichtung. Siehe Abbildung 5.3.

**Definition 5.5** (Skalarprodukt). Seien  $u, v$  zwei Vektoren gleicher Dimension  $n$ . Dann wird das Skalarprodukt berechnet, indem die Vektorelemente komponentenweise multipli-

## 5. Lineare Algebra



**Abbildung 5.3.:** Multiplikation eines Vektors mit  $\lambda_1 = 2$  und  $\lambda_2 = -\frac{1}{2}$ .

ziert und diese Ergebnisse addiert werden. Das Ergebnis des Skalarprodukts ist demnach eine Zahl.

$$u \cdot v = \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_n \end{pmatrix} \cdot \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{pmatrix} := u_1 \cdot v_1 + u_2 \cdot v_2 + \cdots + u_n \cdot v_n = \sum_{i=1}^n u_i v_i$$

Das Skalarprodukt ist *kommutativ* und *assoziativ*.

**Hinweis.**

$$u \cdot v = \cos(\angle(u, v)) \cdot \|u\| \cdot \|v\|$$

So ergibt sich:

- Der Betrag eines Vektors

$$\|v\| = \sqrt{v \cdot v}$$

- Der Kosinus des eingeschlossenen Winkels zweier normierter Vektoren

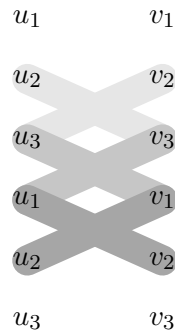
$$\cos(\angle(u, v)) = u \cdot v \quad \text{falls } \|u\| \cdot \|v\| = 1$$

- Die Orthogonalität zweier Vektoren

$$u \perp v \quad \text{falls } u \cdot v = 0$$

**Definition 5.6** (Kreuzprodukt). Seien  $u, v$  zwei Vektoren der Dimension 3. Dann ist das Ergebnis des *Kreuzprodukts* wieder ein Vektor, der wie folgt berechnet wird:

$$\begin{pmatrix} u_1 \\ u_2 \\ u_3 \end{pmatrix} \times \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} := \begin{pmatrix} u_2 v_3 - u_3 v_2 \\ u_3 v_1 - u_1 v_3 \\ u_1 v_2 - u_2 v_1 \end{pmatrix}$$

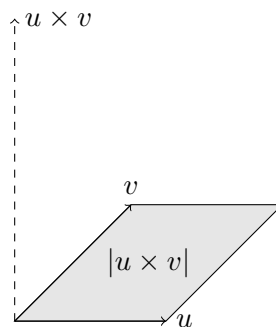


**Abbildung 5.4.:** Diese Skizze erleichtert das Berechnen des Kreuzprodukts. Die Enden einer Linie werden multipliziert. Von dem Produkt der absteigenden Linien wird das Produkt der aufsteigenden subtrahiert. Jedes Kreuz ergibt eine Komponente des Vektorprodukts.

Abbildung 5.4 dient als Merkhilfe. Als weitere Hilfe kann man sich merken, dass in der  $n$ -ten Zeile des Kreuzproduktes keine Elemente aus den  $n$ -ten Zeilen der Vektoren vorkommen.

Das Kreuzprodukt ist *weder* kommutativ noch assoziativ.

**Hinweis.** Das Kreuzprodukt  $u \times v$  ist ein Vektor, welcher *orthogonal* zu  $u$  und  $v$  ist. Der Betrag dieses Vektors entspricht dem Flächeninhalt des Parallelogramms, welches von den Vektoren  $u$  und  $v$  aufgespannt wird, siehe Abbildung 5.5.

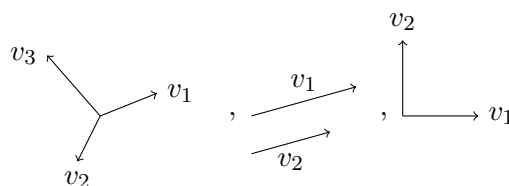


**Abbildung 5.5.:** Das Kreuzprodukt  $u \times v$  zweier Vektoren  $u, v$ .

### 5.1.3. Lineare Abhängigkeit

Lineare Unabhängigkeit ist eine häufig geforderte Eigenschaft einer Menge von Vektoren. So müssen beispielsweise die Basisvektoren, die einen Raum oder eine Ebene aufspannen, immer linear unabhängig sein.

## 5. Lineare Algebra



**Abbildung 5.6.:** Die ersten beiden Mengen von Vektoren sind linear abhängig, die letzte Menge ist linear unabhängig (alle Vektoren im  $\mathbb{R}^2$ ).

**Definition 5.7** (Linearkombination). Eine *Linearkombination* von Vektoren  $v_1, \dots, v_m$  ist eine Addition der Form

$$\sum_{i=1}^m \lambda_i v_i = \lambda_1 v_1 + \dots + \lambda_m v_m$$

wobei  $\lambda_i \in \mathbb{R}$  für  $1 \leq i \leq m$ .  $\lambda_i$  wird auch *Linearfaktor* genannt.

**Definition 5.8** (Lineare Abhängigkeit). Eine Menge von Vektoren  $\{v_1, \dots, v_m\}$  heißt *linear abhängig*, falls es eine Linearkombination gibt mit

$$\sum_{i=1}^m \lambda_i v_i = \mathbf{0}$$

Die *triviale Lösung* ( $\lambda_1 = \lambda_2 = \dots = \lambda_m = 0$ ) ist nicht erlaubt.

Intuitiv ist eine Menge von Vektoren linear abhängig, falls sich mindestens ein Vektor der Menge durch eine Linearkombination anderer Vektoren der Menge darstellen lässt.

**Hinweis.**  $\mathbf{0}$  bezeichnet den *Nullvektor*,  $(0 \ 0 \ \dots \ 0)^T$ .

**Definition 5.9** (Lineare Unabhängigkeit). Eine Menge von Vektoren  $\{v_1, \dots, v_m\}$  heißt *linear unabhängig* genau dann, wenn sie nicht linear abhängig ist.

## 5.2. Matrizen

Unfortunately, no one can be told what The Matrix is. You have to see it for yourself.

*Morpheus*

Matrizen sind neben Vektoren ein weiteres wichtiges Teilgebiet der linearen Algebra. Sie werden beispielsweise verwendet, um die Position und Orientierung von 3D-Objekten in Computerspielen zu codieren. Matrizen gleichen Tabellen, die mit Zahlen und Termen gefüllt sind.

Wir werden zuerst das Konzept einer Matrix sowie verschiedene Matrix-Operationen definieren. Danach werden wir Matrizen auch zum Lösen von sogenannten *linearen Gleichungssystemen* verwenden, welche mit der *linearen Abhängigkeit* eng verbunden sind.

### 5.2.1. Nomenklatur

**Definition 5.10** (Matrix). Eine Matrix  $M$  des Typs  $m \times n$  hat  $m$  Zeilen und  $n$  Spalten.  $a_{ij}$  ist das Element in der Zeile  $i$  und Spalte  $j$

$$\begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}$$

Zeilenvektor  
Spaltenvektor  
Spur

Wir beschäftigen uns ausschließlich mit Matrizen, deren Elemente Zahlen sind, also  $a_{ij} \in \mathbb{R}$  bzw.  $a_{ij} \in \mathbb{C}$ . Für die Menge aller  $m \times n$ -Matrizen über einer Menge  $K$  schreibt man auch  $K^{m \times n}$ , also z. B.  $\mathbb{R}^{m \times n}$ .

### 5.2.2. Matrizenoperationen

**Definition 5.11** (Matrixaddition). Seien  $A, B$  Matrizen des gleichen Typs  $m \times n$ . Dann ist die Addition  $A + B$  ebenfalls eine  $m \times n$ -Matrix. Die beiden Matrizen werden dabei *komponentenweise* addiert, also

$$\begin{aligned} A + B &= \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix} + \begin{pmatrix} b_{11} & b_{12} & \cdots & b_{1n} \\ b_{21} & b_{22} & \cdots & b_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ b_{m1} & b_{m2} & \cdots & b_{mn} \end{pmatrix} \\ &= \begin{pmatrix} a_{11} + b_{11} & a_{12} + b_{12} & \cdots & a_{1n} + b_{1n} \\ a_{21} + b_{21} & a_{22} + b_{22} & \cdots & a_{2n} + b_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} + b_{m1} & a_{m2} + b_{m2} & \cdots & a_{mn} + b_{mn} \end{pmatrix} \end{aligned}$$

Die Matrixaddition ist *assoziativ* und *kommutativ*.

**Definition 5.12** (Multiplikation mit einem Skalar). Sei  $A$  eine Matrix und  $\lambda$  ein Skalar. Dann ergibt sich  $\lambda \cdot A$ , indem jede Komponente von  $A$  mit  $\lambda$  multipliziert wird.

$$\lambda \cdot A = \begin{pmatrix} \lambda \cdot a_{11} & \lambda \cdot a_{12} & \cdots & \lambda \cdot a_{1n} \\ \lambda \cdot a_{21} & \lambda \cdot a_{22} & \cdots & \lambda \cdot a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \lambda \cdot a_{m1} & \lambda \cdot a_{m2} & \cdots & \lambda \cdot a_{mn} \end{pmatrix}$$

**Definition 5.13** (Matrizenmultiplikation). Sei  $A$  eine Matrix des Typs  $m \times n$  und  $B$  eine Matrix des Typs  $n \times k$ . Dann ist  $AB$  eine Matrix des Typs  $m \times k$  und ist wie folgt definiert:

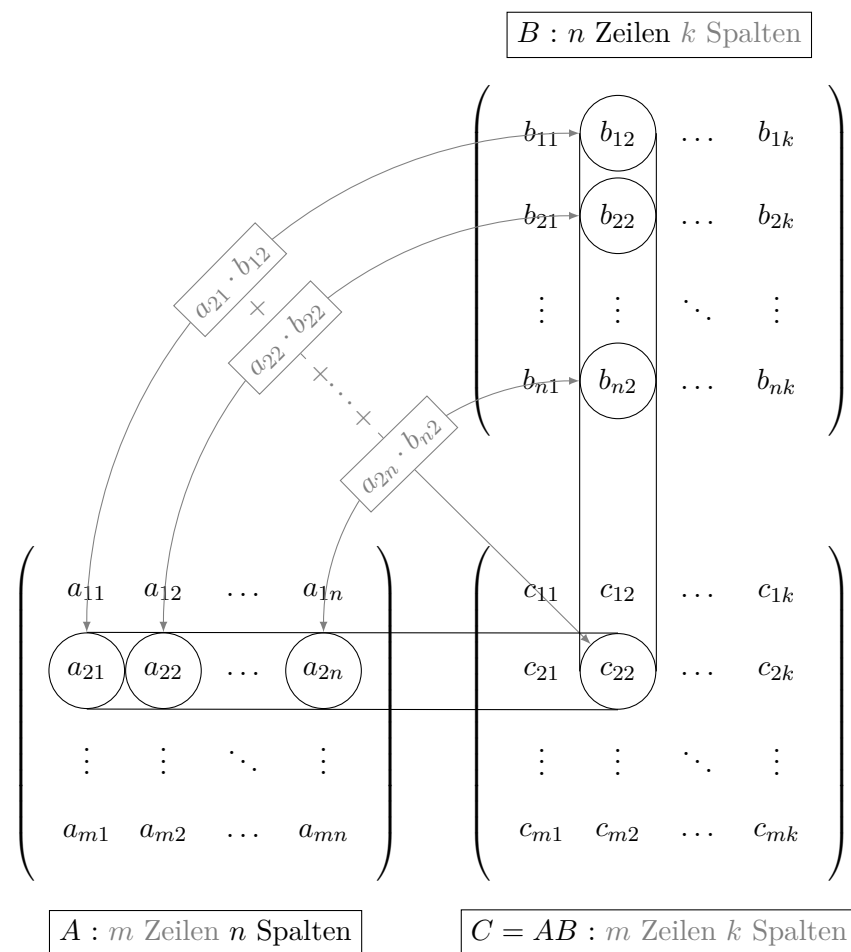
## 5. Lineare Algebra

$$AB = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix} \begin{pmatrix} b_{11} & b_{12} & \dots & b_{1k} \\ b_{21} & b_{22} & \dots & b_{2k} \\ \vdots & \vdots & \ddots & \vdots \\ b_{n1} & b_{n2} & \dots & b_{nk} \end{pmatrix} = \begin{pmatrix} c_{11} & c_{12} & \dots & c_{1k} \\ c_{21} & c_{22} & \dots & c_{2k} \\ \vdots & \vdots & \ddots & \vdots \\ c_{m1} & c_{m2} & \dots & c_{mk} \end{pmatrix}$$

wobei

$$c_{ij} = \sum_{k=1}^n a_{ik} b_{kj}$$

Eine Veranschaulichung bietet Abbildung 5.7.



**Abbildung 5.7.:** Systematische Darstellung einer Matrizenmultiplikation. CC-BY Alain Matthes (<http://www.texample.net/tikz/examples/matrix-multiplication>), angepasste Version.

Die Komponente  $c_{ij}$  ergibt sich also aus dem Skalarprodukt des  $i$ . Zeilenvektors der Matrix  $A$  und des  $j$ . Spaltenvektors der Matrix  $B$ .

**Beispiel 5.14** (Matrizenmultiplikation). Seien die Matrizen  $A, B$  wie folgt gegeben:

$$A = \begin{pmatrix} 1 & 3 & 3 & 7 \\ 4 & -2 & 70 & -45 \end{pmatrix}$$

$$B = \begin{pmatrix} -9 & 50 & -30 \\ -50 & 7 & 42 \\ 6 & -3 & 7 \\ 9 & 1 & -1 \end{pmatrix}$$

Dann können wir  $AB$  bestimmen, denn die Anzahl der Spalten von  $A$  ist gleich der Anzahl der Zeilen von  $B$ :

$$\begin{pmatrix} -9 & 50 & -30 \\ -50 & 7 & 42 \\ 6 & -3 & 7 \\ 9 & 1 & -1 \end{pmatrix}$$

$$\begin{pmatrix} 1 & 3 & 3 & 7 \\ 4 & -2 & 70 & -45 \end{pmatrix} \begin{pmatrix} -78 & 69 & 110 \\ 79 & -69 & 331 \end{pmatrix}$$

Die  $-69$  in Zeile 2 und Spalte 2 erklärt sich beispielsweise folgendermaßen:

$$4 \cdot 50 + (-2) \cdot 7 + 70 \cdot (-3) + (-45) \cdot 1 = -69$$

Die Matrixmultiplikation ist *nicht* kommutativ:

$$\begin{pmatrix} 1 & 0 \\ 2 & 3 \end{pmatrix} \begin{pmatrix} 2 & 3 \\ 1 & 0 \end{pmatrix} = \begin{pmatrix} 1 \cdot 2 + 0 \cdot 1 & 1 \cdot 3 + 0 \cdot 0 \\ 2 \cdot 2 + 3 \cdot 1 & 2 \cdot 3 + 3 \cdot 0 \end{pmatrix} = \begin{pmatrix} 2 & 3 \\ 7 & 6 \end{pmatrix}$$

$$\neq \begin{pmatrix} 8 & 9 \\ 1 & 0 \end{pmatrix} = \begin{pmatrix} 2 \cdot 1 + 3 \cdot 2 & 2 \cdot 0 + 3 \cdot 3 \\ 1 \cdot 1 + 0 \cdot 2 & 1 \cdot 0 + 0 \cdot 3 \end{pmatrix} = \begin{pmatrix} 2 & 3 \\ 1 & 0 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 \\ 2 & 3 \end{pmatrix}$$

### 5.2.3. Berechnung von Determinanten

Jeder quadratischen Matrix  $A \in \mathbb{R}^{n \times n}$  wird ein Skalar zugeordnet, die sogenannte *Determinante*. Man schreibt  $\det A$ . Kennt man die Determinante von  $A$ , so kann man über diese Matrix einige wichtige Aussagen treffen, zum Beispiel ob sie in Stufenform gebracht werden kann (das ihr zugrunde liegende LGS also eindeutig lösbar ist).

Wir wollen an dieser Stelle auf die formale Definition von Determinanten verzichten. Wir werden stattdessen lediglich darauf eingehen, wie sich Determinanten berechnen lassen, da dies häufig benötigt wird, aber viele Studierende zu Beginn damit Probleme haben.

**Definition 5.15** (Determinante einer  $2 \times 2$ -Matrix). Die Determinante einer  $2 \times 2$ -Matrix ist folgendermaßen definiert:

Sei

$$A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

## 5. Lineare Algebra

dann ist

$$\det A = |A| = \begin{vmatrix} a & b \\ c & d \end{vmatrix} := ad - bc$$

Diese Definition ist natürlich nur für  $2 \times 2$ -Matrizen nützlich. Für beliebige  $n \times n$ -Matrizen werden wir den *Laplaceschen Entwicklungssatz*<sup>2</sup> vorstellen, welcher die Berechnung auf  $2 \times 2$ -Matrizen herunterbricht, deren Berechnung wir ja jetzt kennen. Es gibt zudem die *Regel von Sarrus*<sup>3</sup>, die für  $3 \times 3$ -Matrizen ein weniger rechenintensives Verfahren beschreibt.

**Satz 5.16** (Laplacescher Entwicklungssatz). Sei  $A$  eine Matrix vom Typ  $n \times n$  für  $n \geq 3$ . Sei ferner  $A_{ij}$  die Matrix vom Typ  $n-1 \times n-1$ , die aus  $A$  entsteht, indem die  $i$ -te Zeile und die  $j$ -te Spalte gestrichen wird.

Dann ist

$$\det A = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det A_{ij} = \sum_{i=1}^n (-1)^{i+j} a_{ij} \det A_{ij}$$

wobei im ersten Fall  $i$  eine beliebige Zeile und im zweiten Fall  $j$  eine beliebige Spalte der Matrix ist. Man sagt auch, dass man nach der  $i$ -ten Zeile bzw.  $j$ -ten Spalte *entwickelt*.

Anschaulich beschreibt der Laplacesche Entwicklungssatz Folgendes, hier am Beispiel einer beliebigen  $3 \times 3$ -Matrix:

Zu Beginn der Berechnung wird eine beliebige Spalte (bzw. Zeile) gewählt und gestrichen. Im Anschluss wird dann jeweils eine der 3 Zeilen (bzw. Spalten) ebenfalls gestrichen, sodass aus einer  $3 \times 3$ -Matrix drei sogenannte *Streichungsmatrizen* gebildet werden, die jeweils vom Typ  $2 \times 2$  sind. Die Determinanten dieser Matrizen können wir dann berechnen und gemäß der Formel von Laplace zur Determinante der Ausgangsmatrix verrechnen.

**Beispiel 5.17** (Laplacescher Entwicklungssatz). Sei  $A = \begin{pmatrix} 4 & 7 & 9 \\ 6 & 11 & 8 \\ 2 & 4 & 1 \end{pmatrix} \in \mathbb{R}^{3 \times 3}$

Wir entwickeln nach der 1. Zeile, legen also fest, dass  $i = 1$ .

Dann müssen wir folgende Determinanten berechnen:

$$\begin{aligned} |A_{11}| &= \begin{vmatrix} 11 & 8 \\ 4 & 1 \end{vmatrix} = 11 - 32 = -21 \\ |A_{12}| &= \begin{vmatrix} 6 & 8 \\ 2 & 1 \end{vmatrix} = 6 - 16 = -10 \\ |A_{13}| &= \begin{vmatrix} 6 & 11 \\ 2 & 4 \end{vmatrix} = 24 - 22 = 2 \end{aligned}$$

<sup>2</sup>nach Pierre-Simon Laplace, französischer Mathematiker, 1749 - 1827

<sup>3</sup>nach Pierre Frédéric Sarrus, französischer Mathematiker, 1798 - 1861

Also ist nach Satz 5.16

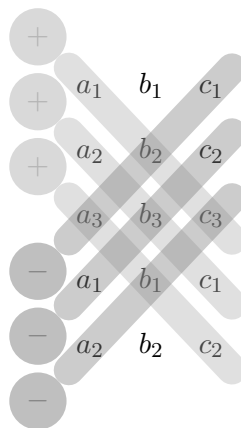
$$\begin{aligned} |A| &= \sum_{j=1}^3 (-1)^{1+j} a_{1j} |A_{1j}| \\ &= (-1)^{1+1} \cdot 4 \cdot -21 + (-1)^{1+2} \cdot 7 \cdot -10 + (-1)^{1+3} \cdot 9 \cdot 2 = 4 \end{aligned}$$

Für komplexere Matrizen muss dieses Verfahren rekursiv angewandt werden. So ergeben sich zum Beispiel für eine  $4 \times 4$ -Matrix bereits  $4 \cdot 3 = 12$   $2 \times 2$ -Matrizen.

**Satz 5.18** (Regel von Sarrus). Für  $3 \times 3$ -Matrizen lässt sich die Determinante auch einfacher mit der *Regel von Sarrus* folgendermaßen berechnen:

$$\det A = \begin{vmatrix} a_1 & b_1 & c_1 \\ a_2 & b_2 & c_2 \\ a_3 & b_3 & c_3 \end{vmatrix} = a_1 b_2 c_3 + a_2 b_3 c_1 + a_3 b_1 c_2 - c_1 b_2 a_3 - c_2 b_3 a_1 - c_3 b_1 a_2$$

Siehe auch Abbildung 5.8.



**Abbildung 5.8.:** Merkbild für die Regel von Sarrus. CC-BY Alain Matthes (<http://www.texample.net/tikz/examples/mnemonic-rule-for-matrix-determinant/>), angepasste Version.

**Satz 5.19** (Eigenschaften von Determinanten). Folgende Eigenschaften von Determinanten werden wir verwenden, ohne jedoch näher auf ihre Herkunft einzugehen bzw. sie zu beweisen.

$\det A \neq 0$  genau dann, wenn

- die Zeilen (Spalten) von  $A$  linear unabhängig sind. Man kann damit also sehr leicht bestimmen, ob  $n$  Vektoren der Dimension  $n$  linear unabhängig sind, indem man sie zu einer Matrix zusammenfasst und ihre Determinante bestimmt.
- $A^{-1}$  existiert,  $A$  ist also *invertierbar*.

### 5.3. Lineare Gleichungssysteme

You shall not pass!

---

*Gandalf, berühmter Zauberer*

Lineare Gleichungssysteme bestehen aus mehreren linearen Gleichungen. Diese enthalten mehrere *Unbekannte*  $x_1, \dots, x_n$ . Durch geschicktes Umstellen der Matrix sollen diese bestimmt werden.

**Definition 5.20** (Lineares Gleichungssystem). Eine *lineare Gleichung* mit  $n$  Unbekannten hat die Form

$$a_1x_1 + a_2x_2 + \dots + a_nx_n = c$$

oder in Vektorschreibweise:  $a \cdot x = c$ , wobei  $c$  eine Konstante und  $a, x$  Vektoren sind.  $x$  heißt auch *Lösungsvektor*.

Ein lineares Gleichungssystem mit  $m$  Gleichungen und  $n$  Unbekannten hat die Form

$$\begin{aligned} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n &= c_1 \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n &= c_2 \\ &\vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n &= c_m \end{aligned}$$

In Vektor- und Matrizenschreibweise kann man ein LGS daher auch formulieren als:

$$A \cdot x = c \text{ oder } \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} c_1 \\ c_2 \\ \vdots \\ c_m \end{pmatrix}$$

Ein lineares Gleichungssystem kann entweder keinen, genau einen oder unendlich viele Lösungsvektoren haben. Damit ein LGS genau einen Lösungsvektor hat, gilt die notwendige aber nicht hinreichende Bedingung, dass es mindestens so viele Gleichungen wie Unbekannte gibt.

**Hinweis.** Ist  $\det A \neq 0$ , so ist  $A \cdot x = c$  eindeutig lösbar und die Lösung ist gegeben durch  $x = A^{-1} \cdot c$ .

**Hinweis.** Da der Lösungsvektor  $x$  sowieso unbekannt ist, schreibt man ein lineares Gleichungssystem meist vereinfachend:

$$\left( \begin{array}{cccc|c} a_{11} & a_{12} & \dots & a_{1n} & c_1 \\ a_{21} & a_{22} & \dots & a_{2n} & c_2 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} & c_m \end{array} \right)$$

**Beispiel 5.21** (Lineares Gleichungssystem). Gegeben sei das folgende Gleichungssystem

$$\begin{aligned}2x_1 + 4x_3 &= 7 \\3x_2 + 3x_3 &= 4 \\2x_1 - 3x_2 + x_3 &= 7\end{aligned}$$

Wir schreiben dies als Matrix wie folgt:

$$\left(\begin{array}{ccc|c}2 & 0 & 4 & 7 \\0 & 3 & 3 & 4 \\2 & -3 & 1 & 7\end{array}\right)$$

### 5.3.1. Substitutionsverfahren

Das wohl intuitivste Verfahren, um ein LGS zu lösen, ist das *Substitutionsverfahren*.

Die Idee ist es, zuerst eine Gleichung nach einer Unbekannten aufzulösen und dann die Unbekannte in den anderen Gleichungen durch den Ausdruck zu ersetzen, der ja nun eine Unbekannte weniger hat. Dies führt man solange weiter durch, bis man den Wert einer der Unbekannten ablesen kann. Durch Rücksubstitution lassen sich dann Lösungen für alle Variablen ausrechnen.

**Beispiel 5.22** (Substitutionsverfahren). Da das grundsätzliche Verfahren recht offensichtlich ist, demonstrieren wir es hier nur mit Hilfe eines LGS mit zwei Unbekannten:

$$\begin{aligned}-2x + 3y &= -8 & \text{I} \\10x + 10y &= 30 & \text{II}\end{aligned}$$

Gleichung I nach x auflösen:

$$x = \frac{1}{2}(3y + 8) = \frac{3}{2}y + 4 \quad \text{III}$$

III in II einsetzen:

$$\begin{aligned}10\left(\frac{3}{2}y + 4\right) + 10y &= 30 \\15y + 40 + 10y - 30 &= 0 \\y &= -\frac{2}{5}\end{aligned}$$

Einsetzen von  $y = -\frac{2}{5}$  in III :  $x = \frac{3}{2}\left(-\frac{2}{5}\right) + 4 = \frac{17}{5}$

### 5.3.2. Gaußsches Eliminationsverfahren

Da das Substitutionsverfahren viel Umformarbeit erfordert, hat Gauß einen Algorithmus, das *Gaußsche Eliminationsverfahren*, entwickelt, mit dessen Hilfe ein LGS nach einem festen Schema gelöst werden kann. Es gliedert sich in zwei Schritte:

## 5. Lineare Algebra

- Umformen des LGS auf *Stufenform*.
- Rücksubstitution der Variablen.

Die *Stufenform* bezeichnet ein LGS, in dem die Spur keine 0 enthält und der Teil unter der Spur nur aus Nullen besteht.

**Beispiel 5.23** (Ein LGS in Stufenform).

$$\left( \begin{array}{ccc|c} 2 & 1 & 0 & 3 \\ 0 & -1 & 3 & 1 \\ 0 & 0 & 1 & 2 \end{array} \right)$$

Um ein LGS auf Stufenform zu bringen, gibt es folgende gültige Operationen:

- i) Multiplikation einer Zeile mit einer Konstanten.
- ii) Addition/Subtraktion einer Zeile mit (dem  $\lambda \neq 0$ -fachen) einer anderen Zeile.
- iii) Vertauschen zweier Zeilen.
- iv) Vertauschen zweier Spalten (hierbei müssen entsprechend auch die Komponenten des Lösungsvektors vertauscht werden).

**Hinweis.** Die letzten beiden Operationen sind optional, sie können jedoch das Umformen vereinfachen.

**Beispiel 5.24** (Gaußsches Eliminationsverfahren).

$$\begin{pmatrix} 2 & 3 & 4 & | & 7 \\ 4 & 12 & 12 & | & 16 \\ 2 & -3 & 1 & | & 7 \end{pmatrix} \xrightarrow{II-2I} \begin{pmatrix} 2 & 3 & 4 & | & 7 \\ 0 & 6 & 4 & | & 2 \\ 2 & -3 & 1 & | & 7 \end{pmatrix} \xrightarrow{III-I} \begin{pmatrix} 2 & 3 & 4 & | & 7 \\ 0 & 6 & 4 & | & 2 \\ 0 & -6 & -3 & | & 0 \end{pmatrix} \xrightarrow{III+II} \begin{pmatrix} 2 & 3 & 4 & | & 7 \\ 0 & 6 & 4 & | & 2 \\ 0 & 0 & 1 & | & 2 \end{pmatrix}$$

Die Matrix ist jetzt in Stufenform und wir wenden Substitution an:

$$\begin{aligned} x_3 &= 2 \\ 6x_2 + 4 \cdot 2 &= 2 \Leftrightarrow x_2 = -1 \\ 2x_1 + 3 \cdot -1 + 4 \cdot 2 &= 7 \Leftrightarrow x_1 = 1 \end{aligned}$$

## 6. Logik und Beweise

“I refuse to prove that I exist,” says God, “for proof denies faith, and without faith I am nothing.”  
“But,” says Man, “the Babel fish is a dead giveaway, isn’t it? It could not have evolved by chance. It proves you exist, and so therefore, by your own arguments, you don’t. QED.”  
“Oh dear,” says God, “I hadn’t thought of that,” and promptly vanishes in a puff of logic.  
“Oh, that was easy,” says Man, and for an encore goes on to prove that black is white and gets himself killed on the next pedestrian crossing.

---

*Douglas Adams, The Hitchhiker’s Guide to the Galaxy*

In der Schule beschäftigt man sich in Mathe hauptsächlich nur mit dem Anwenden von gegebenen Sätzen. An der Universität nehmen die Sätze zusammen mit *Beweisen* eine viel wichtigere Rolle ein. Wichtig ist nicht nur zu wissen, *wie* man einen Satz anwendet, sondern auch zu wissen, *warum* dieser korrekt ist. Nur so kann man eines Tages selbst Aussagen formulieren und beweisen.

### 6.1. Aussagenlogik

In der Aussagenlogik beschäftigt man sich, wie der Name schon verrät, mit Aussagen und deren *Wahrheitswert*. Betrachten wir beispielsweise die Aussage “Pac-Man ist gelb”. Diese Aussage ist offensichtlich *wahr*. Die Aussage “Kühe sind lila” hingegen ist, zur Enttäuschung aller, *falsch*.

Aussagen in der Aussagenlogik sind immer entweder *wahr* oder *falsch*. Die Zuordnung, welche Aussagen wahr oder falsch sind, wird durch die Funktion  $\nu$  festgelegt. Wir wollen im Folgenden all dies in ein formales, mathematisches Modell überführen.

**Definition 6.1** (atomare Aussagen). Sei  $X = \{x_1, x_2, \dots\} = \{x_n \mid n \in \mathbb{N}\}$  eine Menge<sup>1</sup> von sogenannten *Atomen* oder *atomaren Aussagen*.

**Beispiel 6.2.**

- $x_1 := \text{Pac-man ist gelb.}$

---

<sup>1</sup>Wir sind hier davon ausgegangen, dass die einzelnen Elemente mit natürlichen Zahlen durchnummeriert sind. Diese Eigenschaft ist für die Aussagenlogik auch sehr wichtig, denn man fordert, dass  $X$  eine *abzählbare Menge* ist - dies führt an dieser Stelle jedoch zu weit.

## 6. Logik und Beweise

- $x_2 :=$  Kühe sind lila.
- $x_3 :=$  Die Seiten eines Quadrats sind gleich lang.

### 6.1.1. Junktoren

Die Aussagenlogik gibt uns die Möglichkeit, Aussagen durch *Junktoren* zu neuen Aussagen zu kombinieren. So können wir beispielsweise eine Aussage *negieren*.

**Definition 6.3** (Menge der aussagenlogischen Formeln). Die Menge der aussagenlogischen Formeln ist *rekursiv* wie folgt definiert.

- Jede atomare Aussage  $x_i \in X$  ist eine Formel.
- Wenn  $\phi$  eine Formel ist, dann ist auch  $\neg\phi$  eine Formel (Negation).
- Wenn  $\phi$  und  $\psi$  Formeln sind, dann auch
  - $(\phi \wedge \psi)$  (Konjunktion)
  - $(\phi \vee \psi)$  (Disjunktion)
  - $(\phi \rightarrow \psi)$  (Implikation)
  - $(\phi \leftrightarrow \psi)$  (Äquivalenz)

**Beispiel 6.4** (Aussagenlogische Formel). Aus den Formeln  $x_1$  und  $x_2$  lässt sich eine neue Formel  $(x_1 \wedge x_2)$  erstellen. Aus dieser Formel und der Formel  $x_3$  lässt sich eine neue Formel  $((x_1 \wedge x_2) \rightarrow x_3)$  erstellen.

Bislang handelt es sich bei Formeln ausschließlich um *syntaktische Gebilde*. Sie haben noch keine Bedeutung (*Semantik*).

**Definition 6.5** (Wahrheitswerte und Wahrheitsbelegungen). Die Menge der *Wahrheitswerte* ist  $\mathcal{B} = \{0, 1\}$ . 0 steht hierbei für *falsch* und 1 für *wahr*.

Eine *Wahrheitsbelegung*  $\nu$  ordnet jeder atomaren Aussage einen Wahrheitswert zu,  $\nu$  ist also eine Funktion  $\nu : X \rightarrow \mathcal{B}$ . Wir wollen die Wahrheitsbelegung nun in geeigneter Weise erweitern, um beliebigen *Formeln* einen Wahrheitswert zuzuweisen. Die Fortführung einer Wahrheitsbelegung  $\nu$  auf beliebige Formeln bezeichnet man mit  $\hat{\nu}$ . Sie wird ebenfalls *rekursiv* definiert.

Für atomare Aussagen  $x_i$  ist  $\hat{\nu}(x_i) := \nu(x_i)$ . Für alle weiteren Aussagen verwenden wir *Wahrheitstabellen*.

| $\hat{\nu}(\phi)$ | $\hat{\nu}(\neg\phi)$ |
|-------------------|-----------------------|
| 0                 | 1                     |
| 1                 | 0                     |

Die Tabelle ist folgendermaßen zu lesen: Kennen wir bereits den Wahrheitswert für  $\phi$ , so können wir den Wahrheitswert für  $\neg\phi$  aus der Tabelle ablesen.

Die *Negation*  $\neg\phi$  verneint also eine Aussage  $\phi$ .

| $\hat{\nu}(\phi)$ | $\hat{\nu}(\psi)$ | $\hat{\nu}((\phi \wedge \psi))$ |
|-------------------|-------------------|---------------------------------|
| 0                 | 0                 | 0                               |
| 0                 | 1                 | 0                               |
| 1                 | 0                 | 0                               |
| 1                 | 1                 | 1                               |

Die *Konjunktion*  $(\phi \wedge \psi)$  ist eine *Und-Verknüpfung*. Sie ist nur dann wahr, wenn *beide* Aussagen  $\phi$  und  $\psi$  wahr sind (letzte Zeile). Ist auch nur eine der Aussagen falsch, so ist auch die Konjunktion falsch.

| $\hat{\nu}(\phi)$ | $\hat{\nu}(\psi)$ | $\hat{\nu}((\phi \vee \psi))$ |
|-------------------|-------------------|-------------------------------|
| 0                 | 0                 | 0                             |
| 0                 | 1                 | 1                             |
| 1                 | 0                 | 1                             |
| 1                 | 1                 | 1                             |

Die *Disjunktion*  $(\phi \vee \psi)$  ist eine *Oder-Verknüpfung*. Sie ist dann wahr, wenn *mindestens* eine der Aussagen  $\phi$  bzw.  $\psi$  wahr ist. Nur wenn beide Aussagen falsch sind, ist die Disjunktion falsch (erste Zeile).

| $\hat{\nu}(\phi)$ | $\hat{\nu}(\psi)$ | $\hat{\nu}((\phi \rightarrow \psi))$ |
|-------------------|-------------------|--------------------------------------|
| 0                 | 0                 | 1                                    |
| 0                 | 1                 | 1                                    |
| 1                 | 0                 | 0                                    |
| 1                 | 1                 | 1                                    |

Die *Implikation*  $(\phi \rightarrow \psi)$  ist eine *Wenn-Dann-Verknüpfung*. Die Implikation ist die Form der Aussage, die am Anfang wohl am meisten Probleme bereitet, daher erklären wir sie durch ein Beispiel.

**Beispiel 6.6** (Implikation). Für eine natürliche Zahl  $n \in \mathbb{N}$  betrachten wir die Aussagen

1.  $n$  ist durch 2 teilbar.
2.  $n$  ist durch 4 teilbar.

## 6. Logik und Beweise

Es gilt: wenn eine Zahl durch 4 teilbar ist, dann ist sie auch durch 2 teilbar. Es gilt also, dass aus  $n$  ist durch 4 teilbar folgt, dass  $n$  durch 2 teilbar ist. In der Wahrheitstabelle bedeutet dies: wann immer die zweite Aussage wahr ist, ist auch die erste Aussage wahr. Wir können nie in den Fall kommen, dass die zweite Aussage wahr ist, die erste jedoch nicht.

Die *umgekehrte* Folgerung gilt jedoch nicht! Wenn  $n$  durch 2 teilbar ist, muss  $n$  noch nicht durch 4 teilbar sein, es *kann* jedoch auch sein, dass sie auch durch 4 teilbar ist.

Nun zu den (möglicherweise) verwirrenden ersten beiden Zeilen der Wahrheitstabelle. Diese besagen, dass aus einer falschen Annahme eine beliebige Aussage folgt. Man sagt auch *aus Falschem folgt Beliebiges* (Ex falso quodlibet).

In unserem Fall: es kann durchaus sein, dass eine Zahl nicht durch 4 teilbar ist und auch nicht durch 2 (z.B. 5). Es kann aber auch sein, dass eine Zahl nicht durch 4, aber durch 2 teilbar ist (z.B. 6).

Wenn die erste Aussage nicht gilt, dann kann man (auf Basis der Implikation) keine Aussage über den Wahrheitsgehalt der zweiten Aussage machen. Es ist möglich, dass diese Aussage wahr oder falsch ist.

**Beispiel 6.7** (Implikation). Stellen wir uns eine geometrische Form  $A$  vor. Gegeben seien die Aussagen:

- $x_1 := A$  ist ein Quadrat.
- $x_2 := A$  hat vier Ecken.

Es gilt nun offensichtlich, dass, wenn  $x_1$  gilt, auch  $x_2$  gilt. Es kann jedoch auch passieren, dass  $x_2$  gilt, ohne dass  $x_1$  gilt (z.B.  $A$  ist ein Rechteck, aber kein Quadrat).

Daher kann man aus  $x_2$  nicht  $x_1$  folgern.

| $\hat{v}(\phi)$ | $\hat{v}(\psi)$ | $\hat{v}((\phi \leftrightarrow \psi))$ |
|-----------------|-----------------|----------------------------------------|
| 0               | 0               | 1                                      |
| 0               | 1               | 0                                      |
| 1               | 0               | 0                                      |
| 1               | 1               | 1                                      |

Die *Äquivalenz* ( $\phi \leftrightarrow \psi$ ) wird auch als *Genau-Dann-Wenn-Verknüpfung* bezeichnet. Sie ist nur dann wahr, wenn  $\phi$  und  $\psi$  den gleichen Wahrheitswert haben.

Bei der Implikation haben wir gesehen, dass wir bei einer Aussage der Form  $x_1 \rightarrow x_2$  keine Aussage über den Wahrheitswert von  $x_2$  machen können, wenn  $x_1$  falsch ist. Bei der Genau-Dann-Wenn-Aussage ist dies *nicht mehr der Fall*. Wir wissen hier, dass wenn  $x_1$  falsch ist auch  $x_2$  falsch sein muss. Diese Aussage kann man auch folgendermaßen verstehen:  $\phi \leftrightarrow \psi$  ist äquivalent zu  $(\phi \rightarrow \psi) \wedge (\psi \rightarrow \phi)$ . Dieses Vorgehen, die Äquivalenz als Implikation in beide Richtung aufzufassen, wird später noch bei vielen Beweisen wichtig.

**Beispiel 6.8** (Äquivalenz). Stellen wir uns eine geometrische Form  $A$  vor. Gegeben seien die Aussagen:

- $x_1 := A$  ist ein Quadrat.
- $x_2 := A$  hat vier Ecken.
- $x_3 := A$  hat genau vier gleichlange Seiten, die durch rechtwinklige Ecken verbunden sind.

Dann gilt beispielsweise  $x_1 \leftrightarrow x_3$ .

**Beispiel 6.9** (Auswertung). Gegeben seien die Formel  $\phi = ((x_1 \wedge x_2) \rightarrow (x_2 \vee x_3))$  und die Wahrheitsbelegung  $\nu$ , die wie folgt definiert ist:

$$\begin{aligned}\nu(x_1) &= 1 \\ \nu(x_2) &= 0 \\ \nu(x_3) &= 1\end{aligned}$$

Dann können wir  $\hat{\nu}((x_1 \wedge x_2) \rightarrow (x_2 \vee x_3))$  berechnen, indem wir rekursiv die Werte  $\hat{\nu}(x_1 \wedge x_2)$  und  $\hat{\nu}(x_2 \vee x_3)$  berechnen. Wir lesen ab:

$$\begin{aligned}\hat{\nu}(x_1 \wedge x_2) &= 0 \\ \hat{\nu}(x_2 \vee x_3) &= 1\end{aligned}$$

und damit  $\hat{\nu}((x_1 \wedge x_2) \rightarrow (x_2 \vee x_3)) = 1$ .

### 6.1.2. Erfüllbarkeit und Äquivalenzen

Consistency is the last refuge of the unimaginative.

---

*Oscar Wilde*

Nachdem wir nun wissen, was aussagenlogische Formeln sind wollen wir uns mit wichtigen Eigenschaften von solchen Formeln beschäftigen.

**Definition 6.10** (Erfüllbare Formel). Wir definieren die Relation  $\models$  für eine Wahrheitsbelegung  $\nu$  und eine Formel  $\phi$  wie folgt:

$$\nu \models \phi \text{ gdw. } \hat{\nu}(\phi) = 1$$

Eine Formel  $\phi$  heißt *erfüllbar*, wenn es eine Wahrheitsbelegung  $\nu : X \rightarrow \mathcal{B}$  gibt mit  $\nu \models \phi$ . Man nennt  $\nu$  dann auch ein *Modell* für  $\phi$ .

| $\nu(x_1)$ | $\hat{\nu}(\neg x_1)$ | $\hat{\nu}(x_1 \wedge \neg x_1)$ |
|------------|-----------------------|----------------------------------|
| 0          | 1                     | 0                                |
| 1          | 0                     | 0                                |

**Tabelle 6.1.:** Beispiel einer unerfüllbaren Formel. Für keine Belegung der Atome nimmt ihr Wahrheitswert den Wert 1 an.

**Beispiel 6.11** (Erfüllbarkeit). Die Formel  $(x_1 \wedge \neg x_1)$  ist nicht erfüllbar. Es ist nicht möglich, dass  $x_1$  wahr und falsch ist. Um dies zu beweisen, kann man einfach alle möglichen Wahrheitsbelegungen für  $x_1$  durchprobieren. Es gibt zwei Fälle:  $\nu(x_1) = 0$  oder  $\nu(x_1) = 1$ , siehe Tabelle 6.1. Im hinteren Teil der Wahrheitstabelle rechnet man für alle Teilfolgen die entsprechenden Wahrheitswerte aus.

Wir haben gesehen: Für alle möglichen Belegungen von  $x_1$  ist das Ergebnis falsch.

Die Formel  $\phi = ((x_1 \wedge x_2) \wedge x_3)$  hingegen ist erfüllbar, siehe Tabelle 6.2.

| $\nu(x_1)$ | $\nu(x_2)$ | $\nu(x_3)$ | $\hat{\nu}(x_1 \wedge x_2)$ | $\hat{\nu}((x_1 \wedge x_2) \wedge x_3)$ |
|------------|------------|------------|-----------------------------|------------------------------------------|
| 0          | 0          | 0          | 0                           | 0                                        |
| 0          | 0          | 1          | 0                           | 0                                        |
| 0          | 1          | 0          | 0                           | 0                                        |
| 0          | 1          | 1          | 0                           | 0                                        |
| 1          | 0          | 0          | 0                           | 0                                        |
| 1          | 0          | 1          | 0                           | 0                                        |
| 1          | 1          | 0          | 1                           | 0                                        |
| 1          | 1          | 1          | 1                           | 1                                        |

**Tabelle 6.2.:** Beispiel einer erfüllbaren Formel. Es gibt eine Belegung der Atome, sodass ihr Wahrheitswert 1 wird.

Die Belegung  $\nu^*$  mit

$$\nu^* = \{x_1 \rightarrow 1, x_2 \rightarrow 1, x_3 \rightarrow 1\}$$

ist offensichtlich ein Modell für  $\phi$ , also  $\nu^* \models \phi$ .

**Definition 6.12** (Tautologie). Eine Formel  $\phi$  heißt *Tautologie* oder *allgemeingültig*, falls für jede Wahrheitsbelegung  $\nu$  gilt, dass  $\nu \models \phi$ . Egal wie die Variablen der Formel also belegt werden, die Formel ist wahr.

**Beispiel 6.13** (Tautologie). Die Formel  $((x_1 \wedge x_2) \rightarrow x_2)$  ist eine Tautologie wie Tabelle 6.3 zeigt.

| $\nu(x_1)$ | $\nu(x_2)$ | $\hat{\nu}(x_1 \wedge x_2)$ | $\hat{\nu}((x_1 \wedge x_2) \rightarrow x_2)$ |
|------------|------------|-----------------------------|-----------------------------------------------|
| 0          | 0          | 0                           | 1                                             |
| 0          | 1          | 0                           | 1                                             |
| 1          | 0          | 0                           | 1                                             |
| 1          | 1          | 1                           | 1                                             |

**Tabelle 6.3.:** Beispiel einer allgemeingültigen Formel. Für jede beliebige Belegung der Atome ist ihr Wahrheitswert 0.

**Hinweis.** Oftmals verwendet man auch die Zeichen  $\top$  bzw.  $\perp$  als Formeln, wobei für jede Belegung gilt, dass  $\nu \models \top$  und  $\nu \not\models \perp$ .

Es gilt also:

$$\hat{\nu}(\perp) = 0$$

$$\hat{\nu}(\top) = 1$$

$\perp$  und  $\top$  sind somit die syntaktischen Pendanten von 0 und 1.

**Definition 6.14** (Semantische Äquivalenz). Zwei Formeln  $\phi, \psi$  heißen *semantisch äquivalent*, geschrieben  $\phi \equiv \psi$ , falls für alle Belegungen  $\nu$  gilt, dass  $\hat{\nu}(\phi) = \hat{\nu}(\psi)$ .

**Hinweis.** Man sollte das Zeichen  $\equiv$  nicht mit dem Gleichheitszeichen  $=$  verwechseln. Letzteres würde in diesem Zusammenhang bedeuten, dass die Formeln aus der gleichen Zeichenfolge bestehen.

**Beispiel 6.15** (Semantische Äquivalenz). Für die Formeln  $\phi = (x_1 \rightarrow x_2)$  und  $\psi = (\neg x_2 \rightarrow \neg x_1)$  gilt  $\phi \equiv \psi$ . Wir überprüfen alle möglichen Belegungen, siehe Tabelle 6.4.

| $\nu(x_1)$ | $\nu(x_2)$ | $\hat{\nu}(x_1 \rightarrow x_2)$ | $\hat{\nu}(\neg x_2)$ | $\hat{\nu}(\neg x_1)$ | $\hat{\nu}(\neg x_2 \rightarrow \neg x_1)$ |
|------------|------------|----------------------------------|-----------------------|-----------------------|--------------------------------------------|
| 0          | 0          | <b>1</b>                         | 1                     | 1                     | <b>1</b>                                   |
| 0          | 1          | <b>1</b>                         | 0                     | 1                     | <b>1</b>                                   |
| 1          | 0          | <b>0</b>                         | 1                     | 0                     | <b>0</b>                                   |
| 1          | 1          | <b>1</b>                         | 0                     | 0                     | <b>1</b>                                   |

**Tabelle 6.4.:** Beispiel semantischer Äquivalenz. Für alle Belegungen der Atome haben die Formeln den gleichen Wahrheitswert.

## 6. Logik und Beweise

**Hinweis.** Um die Schreibweise von logischen Formeln übersichtlich zu halten, verzichten wir auf bestimmte Klammerungen. Hierbei gilt  $\neg$  vor  $\wedge$  vor  $\vee$ . Weiterhin kann man die äußerste Klammer immer weglassen. Dies bedeutet:

$$x_1 \vee \neg x_2 \wedge x_3 := (x_1 \vee (\neg x_2 \wedge x_3))$$

**Satz 6.16** (Grundlegende Äquivalenzen der Aussagenlogik). In der Aussagenlogik gelten folgende Äquivalenzen.

**Kommutativität**  $\phi \wedge \psi \equiv \psi \wedge \phi$   
 $\phi \vee \psi \equiv \psi \vee \phi$

**Assoziativität**  $(\phi \wedge \psi) \wedge \chi \equiv \phi \wedge (\psi \wedge \chi)$   
 $(\phi \vee \psi) \vee \chi \equiv \phi \vee (\psi \vee \chi)$

**Distributivität**  $\phi \wedge (\psi \vee \chi) \equiv (\phi \wedge \psi) \vee (\phi \wedge \chi)$   
 $\phi \vee (\psi \wedge \chi) \equiv (\phi \vee \psi) \wedge (\phi \vee \chi)$

**Absorption**  $(\phi \wedge \psi) \vee \phi \equiv \phi$   
 $(\phi \vee \psi) \wedge \phi \equiv \phi$

**Auslöschung**  $\phi \wedge (\psi \vee \neg\psi) \equiv \phi$   
 $\phi \vee (\psi \wedge \neg\psi) \equiv \phi$

**Satz 6.17** (Weitere Äquivalenzen). Weitere Äquivalenzen können aus diesen grundlegenden Äquivalenzen hergeleitet werden.

**Tautologie**  $\top \equiv (\phi \vee \neg\phi)$   
 $\perp \equiv (\phi \wedge \neg\phi)$

**Neutrale Elemente**  $\phi \wedge \top \equiv \phi$   
 $\phi \vee \perp \equiv \phi$

**Idempotenz**  $\phi \wedge \phi \equiv \phi$   
 $\phi \vee \phi \equiv \phi$

**Kontraposition**  $\phi \rightarrow \psi \equiv \neg\psi \rightarrow \neg\phi$   
Ein kleines Beispiel: „Mein Hut, der hat drei Ecken, drei Ecken hat mein Hut. Und hätt’ er nicht drei Ecken, so wär’s auch nicht mein Hut.“

**Doppelnegation**  $\neg\neg\phi \equiv \phi$

**De-Morgansche Gesetze**  $\neg(\phi \wedge \psi) \equiv \neg\phi \vee \neg\psi$   
 $\neg(\phi \vee \psi) \equiv \neg\phi \wedge \neg\psi$

**Hinweis.** Die De-Morgansche Gesetze werden verwendet, um Konjunktionen in Disjunktionen zu verwandeln und umgekehrt. Zusammen mit der Doppelnegation kann man so zum Beispiel jegliche Konjunktionen aus einer Formel entfernen, allerdings kann dadurch die Formel deutlich länger und komplizierter werden.

## 6.2. Logik der ersten Stufe

Die Aussagenlogik ist eine verhältnismäßig einfache Logik, jedoch auch in ihrer Aussagekraft sehr eingeschränkt. Betrachte beispielsweise Aussagen wie „Für jede natürliche Zahl  $n$  gilt, dass  $n$  einen Nachfolger besitzt, der auch eine natürliche Zahl ist“ oder „Für jeden Vogel  $v$  gilt, dass  $v$  fliegen kann“ (eine Aussage, die natürlich falsch ist). Diese Aussagen lassen sich zwar als ein Atom in der Aussagenlogik auffassen, sind aber nicht direkt in der Logik formulierbar und miteinander in Beziehung zu bringen. So kommt in beiden Aussagen zwar die Formulierung „für jede(n)“ vor, die Aussagenlogik gibt uns jedoch keine Möglichkeit, dies in der Logik selbst zu formulieren.

Für solche Aussagen gibt es die sogenannte *Logik der ersten Stufe*, auch *Prädikatenlogik* genannt. Diese ist etwas anders aufgebaut als die Aussagenlogik. Wir wollen nicht näher auf Details und eine formale Definition eingehen. Für die meisten Vorlesungen reicht es, wenn man weiß, was die verschiedenen Quantoren bedeuten. Diese Quantoren sind  $\forall$  (der All-Quantor) und  $\exists$  (der Existenz-Quantor).

Die Quantoren werden in Verbindung mit Bedingungen oder Eigenschaften verwendet. Der All-Quantor sagt dabei aus, dass *alle* Elemente einer Menge eine gewisse Eigenschaft erfüllen. Der Existenz-Quantor sagt, dass es *mindestens ein* Element einer Menge gibt, welches die Eigenschaft erfüllt. Die Junktoren der Aussagenlogik können weiterhin verwendet werden, hinzu kommen neben den Quantoren auch Relationssymbole.

**Beispiel 6.18** (Formel der Logik der ersten Stufe). Wir wollen die Aussage formulieren, dass es für jede natürliche Zahl  $n$  eine natürliche Zahl gibt welche größer ist als diese.

$$\forall n \in \mathbb{N} : \exists m \in \mathbb{N} : m > n$$

Zu lesen ist diese Aussage als „Für alle  $n \in \mathbb{N}$  gibt es ein  $m \in \mathbb{N}$ , sodass  $m > n$ .“

Neben diesen beiden hier vorgestellten Logiken gibt es übrigens noch viele andere Logiken, mit denen man in den ersten Semestern aber eher wenig zu tun hat, welche aber gerade im Bereich der Künstlichen Intelligenz oder der Verifikation eingesetzt werden können. So gibt es beispielsweise noch Modallogiken (zeitliche Zusammenhänge), Beschreibungslogiken, Logiken höherer Ordnung (allgemein Logiken  $n$ -ter Ordnung), Fuzzy-Logiken oder Default-Logiken.

## 6.3. Beweise

Since people have started to prove the simplest assertions, many of them proved false.

---

Bertrand Russell

### 6.3.1. Wozu Beweisen?

Beweise sind Folgen von anerkannten Schlüssen, die es uns ermöglichen, die Wahrheit einer Aussage mittels bereits bekannter gültiger Sätze zu belegen. Ausgehend von den Axiomen, welche nie bewiesen werden, können so einfache und schließlich auch komplizierte Aussagen belegt werden.

Das Problem ist: Sollte der Beweis für eine grundlegende Aussage falsch sein, so ist der Beweis für alle daraus bewiesenen Aussagen ebenfalls falsch. Damit erklärt sich, wieso es Mathematiker\*innen so wichtig ist, jegliche Aussagen formal zu beweisen.

Zudem spiegelt sich in Beweisen meist auch das Prinzip der Aussage wieder. Hat man also den Beweis verstanden, hat man auch meist die Aussage selbst verstanden.

**Beispiel 6.19** (Goldbach-Vermutung). Im Jahr 1742 stellte der deutsche Mathematiker Christian Goldbach folgende Vermutung auf:

*“Jede gerade Zahl größer als 2 kann als Summe zweier Primzahlen geschrieben werden.”*

Obwohl die Aussage sehr einfach ist, konnte sie bisher weder bewiesen noch widerlegt werden. Teile der Vermutung mit gewissen Einschränkungen wurden bereits bewiesen, zum Beispiel, dass es mit einer Summe von 5 Primzahlen funktioniert. Auch wurden bereits alle Zahlen bis  $4 \cdot 10^{18}$  getestet.

Weitere bislang ungelöste Probleme sind die vom Clay Mathematics Institute (Cambridge) aufgestellten Millennium-Probleme<sup>2</sup>. Dazu zählt zum Beispiel das in der Informatik bekannte Problem ob  $P = NP$  gilt.

### 6.3.2. Direkter Beweis

Eine häufig verwendete Beweistechnik ist der sogenannte *direkte Beweis*. Bei dieser Art des Beweises versucht man ausgehend von der ursprünglichen Aussage nach und nach *gültige Beweisschritte* anzuwenden, bis man bei seinem Ziel ist.

**Beispiel 6.20** (Direkter Beweis). Die Summe zweier ungerader Zahlen ist eine gerade Zahl.

*Beweis.* Seien  $a$  und  $b$  ungerade Zahlen, also  $a = 2k + 1$  und  $b = 2k' + 1$ .

---

<sup>2</sup><http://www.claymath.org/millennium-problems>

Dann ist

$$\begin{aligned} a + b &= 2k + 1 + 2k' + 1 \\ &= 2k + 2k' + 2 \\ &= 2(k + k' + 1) \end{aligned}$$

also  $a + b = 2n$  und damit eine gerade Zahl (für jeweils  $k, k', n \in \mathbb{N}$ ). □

### 6.3.3. Beweis durch Widerspruch (Reductio ad absurdum)

Bei manchen Aussagen fällt ein direkter Beweis manchmal recht schwer. Hier kann ein Widerspruchsbeweis helfen. Denn während es recht schwer ist, eine Aussage für alle Fälle zu zeigen, reicht ein einziges Beispiel, um sie zu widerlegen. Dies macht man sich zunutze und beweist nicht die Aussage, sondern widerlegt deren exaktes Gegenteil.

**Zu zeigen** Die Aussage, welche bewiesen werden soll.

**Annahme** Hier wird das Gegenteil der Aussage als wahr angenommen.

**Ausführung** Ausgehend von der negierten Annahme wird versucht, einen Widerspruch herzuleiten. Ist die negierte Annahme widerlegt, hat man somit die ursprüngliche Aussage bewiesen.

**Beispiel 6.21** (Satz von Euklid). Es gibt unendlich viele Primzahlen.

*Beweis.*

**Annahme** Es gibt endlich viele Primzahlen.

**Beweisführung** Wenn es endlich viele Primzahlen gibt, so muss es eine definierte endliche Menge aller Primzahlen geben. Sei  $\mathbb{P} = \{p_1, p_2, \dots, p_n\}$  für  $n \in \mathbb{N}$  diese endliche Menge aller Primzahlen. Nun betrachtet man die Zahl  $\lambda = \prod_{i=1}^n p_i$ , also das Produkt aller bekannten Primzahlen. Fügt man nun 1 hinzu, so erhält man die Zahl  $\lambda + 1$ , die durch keine der bekannten Primzahlen teilbar sein kann, denn sie sind alle in  $\lambda$  enthalten. Würde  $\lambda$  mit einer bekannten Primzahl multipliziert werden, würde man nicht  $\lambda + 1$  erhalten. Somit muss  $\lambda + 1$  entweder eine neue Primzahl oder ein Produkt unbekannter Primzahlen sein. Somit muss jedoch die obige Annahme, dass  $\mathbb{P}$  alle Primzahlen enthält, falsch sein. Damit ist im Umkehrschluss bewiesen, dass es unendlich viele Primzahlen gibt. □

**Beispiel 6.22.** Es gibt keine zwei natürlichen Zahlen  $a, b \in \mathbb{N}$  für die gilt, dass  $a^2 - b^2 = 10$ .

## 6. Logik und Beweise

*Beweis.* Angenommen, es gibt zwei Zahlen  $a, b \in \mathbb{N}$  mit  $a^2 - b^2 = 10$ . Dann muss offensichtlich gelten, dass  $a > b$  ist. Setzt man nun  $c := a - b$ , so erhält man

$$\begin{aligned} 10 &= a^2 - b^2 \\ &= (a - b) \cdot (a + b) \\ &= c \cdot ((a + b) + (a + b) - (a + b)) \\ &= c \cdot (a + a - a + b + b - b) \\ &= c \cdot (2b + (a - b)) \\ &= c \cdot (2b + c) \end{aligned}$$

Nun muss gelten, dass  $c$  gerade ist. Denn angenommen,  $c$  wäre ungerade, dann wäre auch  $c + 2b$  ungerade (ungerade Zahl + gerade Zahl = ungerade Zahl). Damit wäre auch  $c \cdot (2b + c)$  ungerade (ungerade Zahl  $\cdot$  ungerade Zahl = ungerade Zahl).

Also muss gelten, dass  $c = 2k$  für eine Zahl  $k \in \mathbb{N}$ . Also

$$10 = 2k \cdot (2b + 2k)$$

also

$$\begin{aligned} 5 &= k \cdot 2(b + k) \\ &= 2(b + k)k \end{aligned}$$

Also  $5 = 2 \cdot n$  für  $n \in \mathbb{N}$ , also ist 5 eine gerade Zahl. Widerspruch. □

### 6.3.4. Vollständige Induktion

Eine wichtige Beweistechnik ist die sogenannte *vollständige Induktion*. Diese wird meistens verwendet, um Aussagen über natürliche Zahlen zu beweisen. Man kann jedoch auch die Idee dieser Beweistechnik auf andere Strukturen übertragen, z. B. aussagenlogische Formeln (*strukturelle Induktion*).

Die natürlichen Zahlen wurden *rekursiv* definiert: 0 ist eine natürliche Zahl und eine natürliche Zahl  $n$  hat genau einen Nachfolger  $n^+$ . Die vollständige Induktion greift diese Definition auf. Wollen wir zeigen, dass eine Aussage  $P(n)$  für alle  $n \in \mathbb{N}$  gilt, so sieht ein Induktionsbeweis folgendermaßen aus:

**Induktionsanfang** Zeige, dass die Aussage für 0 gilt, also  $P(0)$  gilt.

**Induktionsvoraussetzung / Induktionsannahme** Nehme an, dass die Aussage für ein beliebiges  $n$  gilt, also  $P(n)$  gilt.

**Induktionsschritt** Unter dieser Annahme zeige, dass auch  $P(n + 1)$  gilt.

Man zeigt also die Implikation: „Aus  $P(n)$  folgt  $P(n + 1)$ “.

Diese Aussage ist häufig eine Gleichung, z. B. das bekannte Beispiel des *Kleinen Gauß*.

**Satz 6.23** („Der kleine Gauß“). Für alle  $n \in \mathbb{N}$  gilt

$$\sum_{i=0}^n i = \frac{n \cdot (n+1)}{2}$$

*Beweis.* Zuerst zeigen wir, dass die Aussage für 0 gilt:

**Induktionsanfang**

$$\sum_{i=0}^0 i = 0 = \frac{0 \cdot (0+1)}{2}$$

**Induktionsannahme** Es gelte

$$\sum_{i=0}^n i = \frac{n \cdot (n+1)}{2}$$

für ein beliebiges  $n \in \mathbb{N}$ .

**Induktionsschritt** Wir zeigen, dass die Aussage unter der Induktionsannahme auch für  $n+1$  gilt.

Wir müssen also zeigen, dass

$$\sum_{i=0}^{n+1} i = \frac{(n+1) \cdot ((n+1)+1)}{2}$$

$$\begin{aligned} \sum_{i=0}^{n+1} i &= (n+1) + \sum_{i=0}^n i \\ &\stackrel{\text{IV}}{=} (n+1) + \frac{n(n+1)}{2} \\ &= \frac{2(n+1)}{2} + \frac{n(n+1)}{2} \\ &= \frac{(n+1)(2+n)}{2} \\ &= \frac{(n+1)((n+1)+1)}{2} \end{aligned}$$

In der zweiten Zeile haben wir die Induktionsannahme verwendet.

□

Nun wollen wir uns noch kurz überlegen, wieso ein Induktionsbeweis gültig ist.

Als Veranschaulichung kann man sich ein Domino-Spiel vorstellen. Nehmen wir an, wir haben die Steine der Reihe nach aufgestellt. Wollen wir nun zeigen, dass alle Steine umfallen, so reicht es zu zeigen, dass der erste Stein umfällt (entspricht dem Induktionsanfang) und dass, falls ein Stein umfällt (entspricht der Induktionsannahme), dieser auch

## 6. Logik und Beweise

den Stein hinter ihm umwirft (entspricht dem Induktionsschritt). Der erste Stein fällt offensichtlich um, da man dies explizit gezeigt hat. Da nun der erste Stein fällt, fällt nun auch der zweite Stein (durch den Induktionsschritt mit  $n = 0$ ), da der zweite Stein fällt, fällt nun auch der dritte Stein (Induktionsschritt mit  $n = 1$ ). Dies lässt sich nun für alle Steine weiter durchführen.

**Satz 6.24** (Kardinalität von Potenzmengen). Für eine endliche Menge  $A$  ist  $|\mathcal{P}(A)| = 2^{|A|}$ .

*Beweis.*

**Induktionsanfang** Sei  $|A| = 0$ , also  $A = \emptyset$ . Dann ist  $\mathcal{P}(A) = \{\emptyset\}$  und damit  $|\mathcal{P}(A)| = 1 = 2^0 = 2^{|A|}$ .

**Induktionsannahme** Für eine gegebene  $n$ -elementige Menge  $A = \{a_1, \dots, a_n\}$  mit  $n \in \mathbb{N}$  beträgt die Mächtigkeit der Potenzmenge  $|\mathcal{P}(A)| = 2^{|A|}$ .

**Induktionsschritt** Sei  $A$  eine Menge mit  $n + 1$  Elementen. Wir nehmen also an, dass  $A = \{a_1, \dots, a_n, a_{n+1}\}$ .

In  $\mathcal{P}(A)$  gibt es jetzt zwei verschiedene Mengen von Teilmengen:

1. Die Menge  $A_1$  mit jenen Teilmengen, die  $a_{n+1}$  nicht enthalten. Dies sind also gerade die Teilmengen von  $\{a_1, \dots, a_n\}$ , also einer  $n$ -elementigen Teilmenge. Dies sind nach Induktionsvoraussetzung  $2^n$  viele.
2. Die Menge  $A_2$  mit jenen Teilmengen, die  $a_{n+1}$  enthalten. Das sind gerade jene Teilmengen aus der Menge  $A_1$ , in die man zusätzlich noch  $a_{n+1}$  einfügt.

Es gilt, dass  $A_1$  und  $A_2$  disjunkt sind, also dass  $A_1 \cap A_2 = \emptyset$ . Dies ergibt sich dadurch, dass alle Teilmengen in  $A_2$  das Element  $a_{n+1}$  enthalten, welches jene aus  $A_1$  per Definition nicht enthalten. Insgesamt ergibt sich damit

$$|\mathcal{P}(A)| = 2^n + 2^n = 2 \cdot 2^n = 2^{n+1}$$

□

## A. Griechische Buchstaben

| Name    | Griechisch, klein       | Griechisch, groß |
|---------|-------------------------|------------------|
| Alpha   | $\alpha$                | A                |
| Beta    | $\beta$                 | B                |
| Gamma   | $\gamma$                | $\Gamma$         |
| Delta   | $\delta$                | $\Delta$         |
| Epsilon | $\varepsilon, \epsilon$ | E                |
| Zeta    | $\zeta$                 | Z                |
| Eta     | $\eta$                  | H                |
| Theta   | $\vartheta, \theta$     | $\Theta$         |
| Iota    | $\iota$                 | I                |
| Kappa   | $\kappa$                | K                |
| Lambda  | $\lambda$               | $\Lambda$        |
| My      | $\mu$                   | M                |
| Ny      | $\nu$                   | N                |
| Xi      | $\xi$                   | $\Xi$            |
| Omikron | $\omicron$              | O                |
| Pi      | $\pi, \varpi$           | $\Pi$            |
| Rho     | $\rho, \varrho$         | P                |
| Sigma   | $\sigma, \varsigma$     | $\Sigma$         |
| Tau     | $\tau$                  | T                |
| Ypsilon | $\upsilon$              | $\Upsilon$       |
| Phi     | $\varphi, \phi$         | $\Phi$           |
| Chi     | $\chi$                  | X                |
| Psi     | $\psi$                  | $\Psi$           |
| Omega   | $\omega$                | $\Omega$         |



## B. Mathematische Symbole

### B.1. Mengenlehre

| Notation                | Beispiel                                                                                     | Bedeutung                                                                                                                |
|-------------------------|----------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------|
| $\{\dots\}$             | $A = \{1, 2, 3, 4\}$                                                                         | Menge mit den entsprechenden Elementen.                                                                                  |
| $\emptyset$ oder $\{\}$ | $A \cap B = \emptyset$                                                                       | leere Menge                                                                                                              |
| $A \subset B$           | $\mathbb{N} \subset \mathbb{Z}$                                                              | $A$ ist echte Teilmenge von $B$ : $A$ ist Teilmenge und ungleich $B$ .                                                   |
| $A \subseteq B$         | $\{1\} \subseteq \{1, 2\} \subseteq \{2, 1\}$                                                | $A$ ist Teilmenge von $B$ (kann auch gleich $B$ sein).                                                                   |
| $A \supset B$           | $\mathbb{Z} \supset \mathbb{N}$                                                              | $A$ ist echte Obermenge von $B$ .                                                                                        |
| $A \supseteq B$         | $\mathbb{Z} \supseteq \mathbb{N}$                                                            | $A$ ist Obermenge von $B$ .                                                                                              |
| $a \in A$               | $0 \in \mathbb{N}$                                                                           | $a$ ist Element von $A$ : ist in $A$ enthalten.                                                                          |
| $b \notin A$            | $\pi \notin \mathbb{N}$                                                                      | $b$ ist nicht Element von $A$ .                                                                                          |
| $\mathcal{P}(A)$        | $\mathcal{P}(\{1, 2\}) = \{\emptyset, \{1\}, \{2\}, \{1, 2\}\}$                              | Potenzmenge von $A$ : Menge aller Teilmengen von $A$ .                                                                   |
| $A \cup B$              | $\{1, 2\} \cup \{5, 6\} = \{1, 2, 5, 6\}$                                                    | Vereinigung: enthält alle Elemente, die in $A$ oder $B$ enthalten sind.                                                  |
| $A \cap B$              | $\{a, o\} \cap \{a, b, c\} = \{a\}$                                                          | Schnittmenge: enthält alle Elemente, die in $A$ und $B$ enthalten sind.                                                  |
| $A \setminus B$         | $\{1, 2, 3\} \setminus \{3, 4\} = \{1, 2\}$                                                  | Differenz: enthält alle Elemente, die in $A$ aber nicht in $B$ enthalten sind.                                           |
| $A \times B$            | $\{1, 2\} \times \{a, b\}$<br>$= \{(1, a), (1, b), (2, a), (2, b)\}$                         | kartesisches Produkt (Menge aller geordneten Paare):<br>$A \times B := \{(a, b) \mid a \in A, b \in B\}$ .               |
| $A^c$                   | $\{1, 2, 3, 4\}^c = \{5, 6, 7, 8, 9, 10\}$<br>bzgl. $\{x \mid x \in \mathbb{N}, x \leq 10\}$ | Komplement (bzgl. einer Grundmenge $G$ ): enthält alle Elemente, die Element von $G$ welche nicht in $A$ enthalten sind. |
| $ A $                   | $ \mathbb{N}  = \infty$                                                                      | Kardinalität / Mächtigkeit: Anzahl der Elemente.                                                                         |

## B.2. Analysis und Arithmetik

| Notation                              | Beispiel                                                        | Bedeutung                                                                              |
|---------------------------------------|-----------------------------------------------------------------|----------------------------------------------------------------------------------------|
| $f : A \rightarrow B$                 | $f : \mathbb{R} \rightarrow \mathbb{C}$                         | Die Funktion $f$ Elemente aus $A$ (Definitionsmenge) Elemente aus $B$ (Wertemenge) ab. |
| $f : a \mapsto b$                     | $f : 21 \mapsto 42$                                             | Die Funktion $f$ bildet das Element $a$ auf das Element $b$ ab.                        |
| $f(a) = b$                            | $f(21) = 42$                                                    |                                                                                        |
| $f^{-1}$                              | $f^{-1} : \mathbb{C} \rightarrow \mathbb{R}$                    | Umkehrfunktion von $f$ : bildet das Bild von $f$ auf das Urbild ab.                    |
| $f'$                                  | $f(x) = x^2, f'(x) = 2x$                                        | Ableitung von $f$                                                                      |
| $(a_n)_{n \in \mathbb{N}}$            | $a_n = n^2$                                                     | die Folge $a_n$ (Abbildung mit Definitionsmenge $\mathbb{N}$ )                         |
| $\lim_{n \rightarrow \infty} a_n = a$ | $\lim_{n \rightarrow \infty} \frac{1}{n} = 0$                   | Grenzwert: strebt $n$ gegen $\infty$ , so strebt $a_n$ gegen $a$ .                     |
| $\sum_{i=k}^n a_i$                    | $\sum_{i=0}^{100} i = 5050$                                     |                                                                                        |
| $S_n := \sum_{i=0}^n a_i$             | $S_n := \sum_{i=1}^n n^2$                                       | Reihe: Folge der Partialsummen von $a_n$                                               |
| $n!$                                  | $4! = 24$                                                       | Fakultät von $n$ : $0! = 1 \quad n! = n \cdot (n-1)!$                                  |
| $\Delta$                              | $\frac{\Delta y}{\Delta x} = \frac{f(x_2) - f(x_1)}{x_2 - x_1}$ |                                                                                        |
| $x = y$                               | $20 + 22 = 42$                                                  | Gleichheit                                                                             |
| $x \neq y$                            | $21 \neq 42$                                                    | Ungleichheit                                                                           |
| $\approx$                             | $\pi \approx 3.14$                                              | Rundung                                                                                |
| $\hat{=}$                             | 14 Punkte $\hat{=}$ Note 1                                      | Entspricht-Zeichen (für Größen unterschiedlicher Art)                                  |
| $ a $                                 | $ -2  = 2$                                                      | Betrag von $a$                                                                         |
| $\lfloor a \rfloor$                   | $\lfloor \pi \rfloor = 3$                                       | größte ganze Zahl $\leq a$                                                             |
| $\lceil a \rceil$                     | $\lceil \pi \rceil = 4$                                         | kleinste ganze Zahl $\geq a$                                                           |
| $\Re$                                 | $\Re(i) = 0$                                                    | Realteil der komplexen Zahl                                                            |
| $\Im$                                 | $\Im(1 + 2i) = 2$                                               | Imaginärteil der komplexen Zahl                                                        |
| $\bar{z}$                             | $2 + i = \overline{2 - i}$                                      | konjugierte komplexe Zahl von $z$                                                      |
| $\prod_{i=k}^n a_i$                   | $\prod_{i=1}^n i = n!$                                          | das Produkt aller Folgenglieder $a_k$ bis $a_n$                                        |
| $[a; b]$                              | $[0; 5]$                                                        | geschlossenes Intervall: $\{x \in \mathbb{R} \mid a \leq x \leq b\}$                   |
| $]a; b[$ bzw. $(a; b)$                | $]0; 5[$                                                        |                                                                                        |
| $\int_a^b f(x) dx$                    | $\int_{-\pi}^{\pi} \sin x dx = 0$                               | das Integral von $a$ bis $b$ über $f(x)$                                               |

## B.3. Vektorrechnung

| Notation     | Beispiel                                                                                                                                      | Bedeutung                      |
|--------------|-----------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------|
| $u \cdot v$  | $\begin{pmatrix} 1 \\ 0 \end{pmatrix} \cdot \begin{pmatrix} 2 \\ 3 \end{pmatrix} = 2$                                                         | Skalarprodukt von $v$ und $w$  |
| $u \times v$ | $\begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} \times \begin{pmatrix} -7 \\ 8 \\ 9 \end{pmatrix} = \begin{pmatrix} -6 \\ -30 \\ 22 \end{pmatrix}$ | Kreuzprodukt von $u$ und $v$   |
| $\ v\ $      | $\left\  \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right\  = \sqrt{2}$                                                                            | Betrag / Länge des Vektors $v$ |

## B.4. Logik

| Notation          | Beispiel                                 | Bedeutung                                          |
|-------------------|------------------------------------------|----------------------------------------------------|
| $\neg$            | $\neg\phi$                               | Negation (nicht $\phi$ )                           |
| $\wedge$          | $\phi \wedge \psi$                       | Konjunktion ( $\phi$ und $\psi$ )                  |
| $\vee$            | $\phi \vee \psi$                         | Disjunktion ( $\phi$ oder $\psi$ )                 |
| $\rightarrow$     | $\phi \rightarrow \psi$                  | Implikation (aus $\phi$ folgt $\psi$ )             |
| $\leftrightarrow$ | $\phi \leftrightarrow \psi$              | Logische Äquivalenz (beidseitige Implikation)      |
| $\equiv$          | $\neg\neg x_1 \equiv x_1$                | Zwei Formeln sind semantisch äquivalent            |
| $\forall x$       | $\forall x \in \mathbb{N} : x \geq 0$    | Allquantor (für alle Elemente gilt)                |
| $\exists x$       | $\exists x \in \mathbb{N} : x + 21 = 42$ | Existenzquantor (es gibt ein Element, sodass gilt) |