

1

Grundlagen

Das grundlegende Handwerkszeug des Mathematikers ist die *Aussagenlogik*, und zu den grundlegenden Begriffen der modernen Mathematik gehören die Begriffe der *Menge* und der *Abbildung*.

Grundlagenfragen sind allerdings immer auch schwierige Fragen. Sie setzen eine tiefe Kenntnis der Materie voraus, um ihre Bedeutung und ihren Reiz zu erschließen. Für einen Neuling sind sie dagegen ein trockenes, wenig appetitanregendes Brot.

Wir werden die Aussagenlogik sowie die Begriffe der Menge und der Abbildung daher nur in einem ›naiven‹ Sinn kurz ansprechen. Wir verzichten auf eine mathematisch präzise Formulierung, da Aufwand und Ertrag für unsere Zwecke in keinem vernünftigen Verhältnis stehen. Im Wesentlichen geht es darum, sich auf einen Sprachgebrauch für alles Weitere zu verständigen.

1.1

Aussagenlogik

Das Gebäude der Mathematik wird mit den Regeln der Aussagenlogik errichtet. Dessen Bausteine sind Aussagen im Sinne von Aristoteles, denen genau einer von zwei möglichen Wahrheitswerten zukommt.

Aristotelischer Aussagebegriff Eine *Aussage* ist entweder *wahr* oder *falsch* – *tertium non datur*¹. ✕

Es gibt in der Mathematik also kein *vielleicht* – solche Aussagen werden gar nicht erst zugelassen.

¹ Ein Drittes gibt es nicht.

► **Beispiele** Aristotelische Aussagen sind:

- A. Eine Woche hat sieben Tage.
- B. 2 ist eine Primzahl.²
- C. Es gibt unendlich viele Primzahlen.
- D. Es gibt unendlich viele Primzahlzwillinge.³

Keine Aussagen im Sinne der Mathematik sind dagegen:

- E. Hoffentlich ist die Vorlesung bald zu Ende.
- F. Schere, Stein, Papier.
- G. Nachts ist es kälter als draußen. ◀

Dabei ist unerheblich, ob der Wahrheitswert einer Aussage *tatsächlich bekannt* ist. So ist *Es gibt unendlich viele Primzahlzwillinge* sicher wahr oder falsch, nur wissen wir dies im Jahr 2018 nicht.

■ Logische Verknüpfungen

Aus Aussagen lassen sich durch logische Verknüpfungen neue Aussagen bilden. Das Ergebnis einer solchen Verknüpfung legt man in einer *Wahrheitstafel* fest. So wird die Verknüpfung zweier Aussagen p und q ⁴ durch *und* und *oder*, in Symbolen $p \wedge q$ respektive $p \vee q$, durch folgende Wahrheitstafel definiert. Dabei stehe 1 für wahr und 0 für falsch⁵.

p	q	$p \wedge q$	$p \vee q$
1	1	1	1
1	0	0	1
0	1	0	1
0	0	0	0

Das logische *und* entspricht dem allgemeinen Sprachgebrauch: $p \wedge q$ ist dann und nur dann wahr, wenn sowohl p als auch q wahr sind.

Beim logischen *oder* ist zu beachten, dass es *einschließend* ist: $p \vee q$ ist wahr, wenn p oder q oder *beide* wahr sind. Das umgangssprachliche *oder* ist dagegen meist *ausschließend* gemeint. *Es regnet oder ich gehe spazieren* meint üblicherweise *Entweder es regnet, oder ich gehe spazieren*.

² Eine natürliche Zahl heißt *prim*, wenn sie genau zwei Teiler besitzt, nämlich 1 und sich selbst. Somit ist 2 eine Primzahl, nicht aber 1.

³ Ein *Primzahlzwillig* liegt vor, wenn p und $p + 2$ prim sind. Zum Beispiel sind 5 und 7 oder 11 und 13 Primzahlzwillinge.

⁴ p und q stehen hier also für irgendwelche Aussagen. Ebenso gut könnte man A und B schreiben, oder α und β , oder beliebige andere Symbole.

⁵ Dies soll jetzt keine Diskriminierung der 0 darstellen.

► **Beispiele** Folgende Aussagen sind wahr:

- A. $0 < 1$ und $\sqrt{2}$ ist irrational.
- B. $0 < 1$ oder $\sqrt{2}$ ist irrational.
- C. $0 < 1$ oder $\sqrt{2}$ ist rational. ◀

Die Negation einer Aussage p durch *nicht*, symbolisch $\neg p$, ändert schlicht den jeweiligen Wahrheitswert in sein Gegenteil:

p	$\neg p$
1	0
0	1

Allerdings kann in konkreten Fällen die korrekte Verneinung einer Aussage durchaus Schwierigkeiten bereiten.

- **Beispiele**
- A. Es ist nicht so, dass, wenn es heute schneit, morgen die Sonne scheint.
 - B. Es gilt nicht $\lim_{t \rightarrow \infty} \sin t = 0$. ◀

Wichtig ist auch die logische Verknüpfung zweier Aussagen p und q durch *wenn-dann* und *genau-dann-wenn*, in Symbolen $p \rightarrow q$ und $p \leftrightarrow q$:

p	q	$p \rightarrow q$	$p \leftrightarrow q$
1	1	1	1
1	0	0	0
0	1	1	0
0	0	1	1

Das logische *genau-dann-wenn* folgt dabei wiederum der Umgangssprache: $p \leftrightarrow q$ ist wahr, wenn p und q denselben Wahrheitsgehalt haben.

Ungewohnt erscheint anfangs dagegen das logische *wenn-dann*: die Verknüpfung $p \rightarrow q$ ist wahr, wenn entweder p und q wahr sind, oder wenn p falsch ist. In letzterem Fall spielt also der Wahrheitsgehalt von q keine Rolle. Egal was ich aus einer falschen Aussage p folgere, die *logische Verknüpfung als Ganzes* wird als wahr gewertet – *ex falso quodlibet*⁶.

► **Beispiel** Die Aussage

$$x > 0 \rightarrow x + 1 > 0$$

⁶ Aus Falschem folgt alles Beliebige.

ist immer wahr. Denn wenn $x > 0$ gilt, so gilt auch $x + 1 > 0$. Ist dagegen $x \leq 0$, so ist die wenn-dann-Verknüpfung nach dieser Konvention ebenfalls wahr, denn diesen Fall wollen wir ja auch gar nicht betrachten. $\blacktriangleleft$

Bemerkung Es gibt insgesamt $2^4 = 16$ verschiedene Möglichkeiten, das Ergebnis einer logischen Verknüpfung zweier Aussagen durch eine Wahrheitstafel festzulegen. Es gibt also genau 16 verschiedene *binäre logische Funktionen*. Diese können alle durch $\wedge$, $\vee$ und $\neg$ dargestellt werden. Tatsächlich reicht sogar eine einzige Funktion, entweder die *nicht-und* oder *nand-Funktion*

$$p \sqcap q :\Leftrightarrow \neg(p \wedge q),$$

oder die *nicht-oder* oder *nor-Funktion*

$$p \sqcup q :\Leftrightarrow \neg(p \vee q).$$

Diese spielen daher beim Aufbau logischer Schaltkreise eine große Rolle. $\rightarrow$

Mit den genannten Verknüpfungen lassen sich sehr komplexe Ausdrücke bilden. Um dabei Klammern zu sparen, legt man fest, dass $\neg$ vor $\vee$, $\wedge$ vor $\rightarrow$ vor $\leftrightarrow$ bindet. So ist $(p \vee (\neg q)) \rightarrow p$ gleichbedeutend mit $p \vee \neg q \rightarrow p$.

■ Tautologien

Ein mathematisches Gebäude wird errichtet, indem man aus gewissen Grundannahmen – den *Axiomen* – und bereits bewiesenen Aussagen durch korrekte logische Schlüsse neue Aussagen gewinnt. Jeder einzelne logische Schluss von einer Aussage p auf eine Aussage q stellt dabei sicher, dass q immer dann wahr ist, wenn p wahr ist. Ist p dagegen falsch, so interessiert q nicht weiter. Mit anderen Worten, die Aussageverknüpfung $p \rightarrow q$ ist *immer wahr*, egal welche Wahrheitswerte p und q annehmen.

Definition Eine *Tautologie* ist eine Aussage, die immer wahr ist. Eine *Kontradiktion* oder ein *Widerspruch* ist eine Aussage, die immer falsch ist. $\times$

► *Beispiele* A. Die einfachste Tautologie ist $p \vee \neg p$.

B. Die einfachste Kontradiktion ist $p \wedge \neg p$.

C. $x > 0 \rightarrow x + 1 > 0$ ist immer wahr, also eine Tautologie. $\blacktriangleleft$

Ist $p \rightarrow q$ immer wahr, so schreibt man

$$p \Rightarrow q,$$

gelesen p *impliziert* q . Genau genommen ist dies eine *Metaaussage*, also eine Aussage über eine Aussage. Mit $p \Rightarrow q$ sagen wir aus, dass wann immer p gilt, dann auch q gilt.

► *Beispiel* $x > 0 \Rightarrow x + 1 > 0$. ◀

Entsprechend steht

$$p \Leftrightarrow q,$$

gelesen p *äquivalent* q , für die Aussage, dass $p \leftrightarrow q$ immer wahr ist. In diesem Fall nehmen p und q immer denselben Wahrheitswert an und stellen somit dieselbe logische Funktion dar. Die folgende Äquivalenz ist zum Beispiel sehr nützlich.

1 **Notiz** $p \rightarrow q \Leftrightarrow \neg p \vee q$. ✕

»»» *Beweis* Dies verifiziert man wieder anhand einer Wahrheitstafel:

p	q	$\neg p$	$p \rightarrow q$	$\neg p \vee q$
1	1	0	1	1
1	0	0	0	0
0	1	1	1	1
0	0	1	1	1

. »»»

Auf dieselbe Weise verifiziert man eine Fülle weiterer logischer Äquivalenzen. Hier sind die wichtigsten betreffend $\wedge$ und $\vee$.

2 **Logische Äquivalenzen** Es gelten die *Distributivgesetze*

$$(p \wedge q) \vee r \Leftrightarrow (p \vee r) \wedge (q \vee r),$$

$$(p \vee q) \wedge r \Leftrightarrow (p \wedge r) \vee (q \wedge r),$$

die *Verschmelzungsgesetze*

$$(p \wedge q) \vee p \Leftrightarrow p \Leftrightarrow (p \vee q) \wedge p,$$

die *Regeln von de Morgan*

$$\neg(p \wedge q) \Leftrightarrow \neg p \vee \neg q,$$

$$\neg(p \vee q) \Leftrightarrow \neg p \wedge \neg q,$$

sowie die *Regel von der doppelten Negation*

$$\neg(\neg p) \Leftrightarrow p. \quad \times$$

Eine $\wedge$ - respektive $\vee$ -Verknüpfung wird also verneint, indem die einzelnen Aussagen verneint und die Junktoren $\wedge$ und $\vee$ vertauscht werden.

► *Beispiele* A. Es ist nicht so, dass, wenn es heute schneit, morgen die Sonne scheint ist gleichbedeutend mit Es schneit, und die Sonne scheint nicht:

$$\neg(\text{snow} \rightarrow \text{sun}) \Leftrightarrow \neg(\neg \text{snow} \vee \text{sun}) \Leftrightarrow \text{snow} \wedge \neg \text{sun}.$$

B. *Es stimmt nicht, dass es regnet oder ich spazieren gehe* ist gleichbedeutend mit *Es regnet nicht, und ich gehe nicht spazieren*:

$$\neg(r \vee s) \Leftrightarrow \neg r \wedge \neg s.$$

$$\text{C. } \neg(0 < 1 \vee \sqrt{2} \in \mathbb{Q}) \Leftrightarrow 0 \geq 1 \wedge \sqrt{2} \notin \mathbb{Q}. \quad \blacktriangleleft$$

■ Beweistechniken

Die nächsten Äquivalenzen bilden die Grundlage für einige der wichtigsten Beweistechniken.

3 Abtrennungs- und Syllogismusregel

$$p \wedge (p \rightarrow q) \Rightarrow q \quad \text{und} \quad (p \rightarrow q) \wedge (q \rightarrow r) \Rightarrow p \rightarrow r. \quad \times$$

⟦ Um die Gültigkeit der ersten Implikation zu beweisen, müssen wir nur den Fall betrachten, wo die linke Seite wahr ist. In diesem sind also p und $p \rightarrow q$ wahr. Das aber ist nur möglich, wenn auch q wahr ist. Also ist auch q wahr.

Entsprechend argumentiert man bei der zweiten Implikation. Aber natürlich kann man beide Aussagen auch mit Wahrheitstafeln beweisen A_3 . ⟧

Diese Regeln beschreiben die Technik des direkten Beweises. Die Abtrennungsregel besagt: *Gilt p , und folgt q aus p , so gilt auch q .* Die Syllogismusregel besagt: *Folgt q aus p , und r aus q , so folgt auch r aus p .* Dies entspricht dem Alltagsgebrauch und soll nicht weiter illustriert werden.

4 Kontrapositionsregel

$$p \rightarrow q \Leftrightarrow \neg q \rightarrow \neg p. \quad \times$$

⟦ Beweis Dies können wir bereits mit den vorhandenen Mitteln formal beweisen, ohne Rückgriff auf eine Wahrheitstafel:

$$\begin{aligned} p \rightarrow q &\Leftrightarrow (\neg p) \vee q \\ &\Leftrightarrow \neg(\neg q) \vee (\neg p) \\ &\Leftrightarrow (\neg q) \rightarrow (\neg p). \quad \rangle\rangle\rangle \end{aligned}$$

Die Kontrapositionsregel bildet die Grundlage des *indirekten Beweises*. Statt q aus p zu folgern, zeigt man, dass die Verneinung von q zur Verneinung von p führt.

- 5 ▶ *Beispiel eines indirekten Beweises* Ist das Quadrat einer natürlichen Zahl n gerade, so ist auch n selbst gerade.

Der Beweis erfolgt indirekt. Wir negieren die Folgerung und nehmen an, dass n *nicht gerade* ist. Dann ist n ungerade und damit von der Form

$$n = 2m + 1$$

mit einer ganzen Zahl $m \geq 0$. Dann aber ist auch

$$n^2 = (2m + 1)^2 = 4m^2 + 4m + 1 = 4(m^2 + m) + 1 = 4k + 1$$

mit der ganzen Zahl $k = m^2 + m \geq 0$. Also ist n^2 ebenfalls ungerade. Dies ist die Negation der Voraussetzung, und der indirekte Beweis ist abgeschlossen. Damit ist auch die ursprüngliche Aussage bewiesen. ◀

6 Widerspruchregel

$$(\neg p \rightarrow q) \wedge (\neg q) \Rightarrow p. \quad \times$$

◀◀◀ Wieder müssen wir nur den Fall betrachten, wo die linke Seite wahr ist. Dann sind $\neg p \rightarrow q$ und $\neg q$ wahr. Also ist q falsch. Dann kann aber $\neg p$ nicht falsch sein, sonst wäre ja $\neg p \rightarrow q$ falsch. Also ist p wahr. ▶▶▶

Die Widerspruchregel ist die Grundlage des *Widerspruchsbeweises*. Um eine Aussage p zu beweisen, nehmen wir an, dass sie *nicht gilt*, also $\neg p$ wahr ist. Können wir daraus einen Widerspruch ableiten, also eine Aussage, die immer falsch ist, so folgt, dass p wahr ist.

7 ▶ Beispiel eines Widerspruchsbeweises $\sqrt{2}$ ist keine rationale Zahl.

Wir nehmen an, $\sqrt{2}$ ist *rational*. Dann ist also

$$\sqrt{2} = r/s$$

mit natürlichen Zahlen r und $s \neq 0$. Wir können annehmen, dass r und s nicht beide gerade sind, denn andernfalls dividieren wir r und s so lange durch 2, bis dieser Zustand erreicht ist.

Aus $\sqrt{2} = r/s$ folgt nun durch Quadrieren

$$2s^2 = r^2.$$

Also ist r^2 gerade. Dann ist auch r selbst gerade, also $r = 2t$ mit einer anderen natürlichen Zahl t . Also gilt $2s^2 = r^2 = (2t)^2 = 4t^2$, oder

$$s^2 = 2t^2.$$

Also ist s ebenfalls gerade.

Somit sind r und s beide gerade, im Widerspruch zur Annahme, dass r und s *nicht beide* gerade sind. Damit haben wir also die Annahme, dass $\sqrt{2}$ rational ist, zu einem Widerspruch geführt. ◀

8 Äquivalenzregel

$$p \leftrightarrow q \Leftrightarrow (p \rightarrow q) \wedge (q \rightarrow p). \quad \times$$

Die Äquivalenz zweier Aussagen ist also gleichbedeutend damit, dass jede der Aussagen aus der jeweils anderen folgt. Salopp gesagt: »um $p \leftrightarrow q$ zu zeigen, muss man beide Richtungen zeigen«.

► *Beispiel* Das Quadrat einer natürlichen Zahl n ist gerade genau dann, wenn n gerade ist.

Den $\rightarrow$ -Teil haben wir oben gezeigt. Der $\leftarrow$ -Teil besteht darin zu zeigen, dass das Quadrat einer geraden natürlichen Zahl ebenfalls gerade ist. Dies ist aber eine leichte Übung. ◀

Es gibt noch eine Reihe weiterer elementarer Beweistechniken, beispielsweise den Induktionsbeweis. Diese werden wir später kennenlernen.

Bemerkung zum mathematischen Sprachgebrauch Gilt $p \Rightarrow q$, so nennt man p *hinreichend für q* , und q *notwendig für p* .

Denn immer wenn p gilt, so gilt auch q . Somit »reicht p aus, damit auch q gilt«. Gilt dagegen q nicht, so gilt aufgrund der Kontrapositionsregel auch p nicht. Somit ist » q notwendig, damit auch p gilt«.

Gilt sogar $p \Leftrightarrow q$, so ist p *hinreichend und notwendig für q* . Und natürlich gilt dies auch umgekehrt. $\rightarrow$

► *Beispiel*

$$x > 4 \Rightarrow x > 2.$$

Und tatsächlich ist $x > 4$ hinreichend dafür, dass auch $x > 2$. Allerdings ist es nicht notwendig. Und umgekehrt ist $x > 2$ notwendig dafür, dass auch $x > 4$ gilt. Allerdings ist es nicht hinreichend. ◀

1.2 Mengen

⁷ Der Begriff der Menge wurde am Ende des 19. Jahrhunderts von Georg Cantor wie folgt eingeführt.

Definition (Cantor 1895) Eine *Menge* ist eine Zusammenfassung M von bestimmten, wohlunterschiedenen Objekten unserer Anschauung oder unseres Denkens (welche die *Elemente* von M genannt werden) zu einem Ganzen. ✕

›So stelle ich mir eine Menge auch vor‹, würde man wohl sagen. Tatsächlich ist dies aber alles andere als eine präzise Definition. Was ist denn ein ›bestimmtes Objekt‹, und aus welchen Objekten besteht unsere ›Anschauung‹ insgesamt?

Es dauerte auch nicht lange, bis die Widersprüchlichkeit dieser Definition erkannt wurde. Bertrand Russell entdeckte um 1901 ein besonders einfaches Beispiel. Bildet man die Menge M aller Mengen X , die sich nicht selbst enthalten – in Symbolen

$$M = \{X : X \notin X\}$$

– so lässt sich nicht entscheiden, ob M sich selbst enthält oder nicht. Nimmt man an, dass $M \in M$, so folgt hieraus $M \notin M$. Nimmt man dagegen an, dass $M \notin M$, so folgt hieraus wiederum $M \in M$ A-12.

Ein analoges, nicht-mathematisches Beispiel ist das Paradoxon des Barbiers eines Dorfes, der alle Männer im Dorf rasiert, die sich nicht selbst rasieren. Auch hier kann man nicht entscheiden, ob der Barbier sich selbst rasiert oder nicht.

Mathematisch ist dieser Zustand natürlich nicht haltbar und führte zur Entwicklung der *axiomatischen Mengenlehre* von Zermelo, Fraenkel und anderen. Hier lassen sich nicht mehr völlig beliebige Objekte zu Mengen zusammenfassen. Insbesondere ist die Russellsche Konstruktion keine Menge mehr, und die Frage, ob M zu M gehört, kann nicht mehr gestellt werden.

Wir werden uns aber nicht mit der axiomatischen Mengenlehre beschäftigen – so wie es auch die meisten Mathematiker halten. Alle Mengen, die wir im naiven Sinne bilden, sind Mengen auch im axiomatischen Sinn, und das soll uns genügen.

■ Mengen

Für unseren Gebrauch ist eine Menge bestimmt durch die in ihr enthaltenen Elemente. Ist M eine Menge, so ist ein beliebiges Objekt m – wieder so ein

⁷ Dieses Thema wird in der Vorlesung zur Linearen Algebra behandelt.

unbestimmter Begriff – entweder *Element* von M , geschrieben

$$m \in M,$$

oder *nicht Element* von M , geschrieben

$$m \notin M.$$

Es ist also

$$m \notin M \Leftrightarrow \neg(m \in M).$$

Insbesondere enthält eine Menge jedes ihrer Elemente nur einmal, und auf deren Reihenfolge kommt es nicht an. Zwei Mengen sind *gleich*, wenn sie dieselben Elemente enthalten.

Es gibt auch eine Menge ohne Elemente, die sogenannte *leere Menge*

$$\emptyset = \{\}.$$

Für jedes beliebige Objekt m gilt also $m \notin \emptyset$.

9 Notiz Es gibt nur eine leere Menge. ✕

⟦⟦⟦⟦ *Beweis* Seien $\emptyset_1$ und $\emptyset_2$ zwei leere Mengen. Dann ist jedes Element, das in $\emptyset_1$ enthalten ist, auch in $\emptyset_2$ enthalten. Das Umgekehrte gilt ebenfalls. Also enthalten beide Mengen dieselben Elemente und sind damit gleich. ⟧⟧⟧⟧

Enthält eine Menge nur endlich viele Elemente, so können wir sie – jedenfalls im Prinzip – durch Aufzählung aller ihrer Elemente angeben:

$$\{\text{H, i, l, f, e}\}, \quad \{\boxminus, \boxtimes, \boxplus\}.$$

Bei der Aufzählung ist es erlaubt, Elemente einer Menge mehrfach zu *nennen*, auch wenn sie nur einmal *enthalten* sind. Dies ist eine praktische Konvention.

► Beispiele

$$\begin{aligned} \boxplus &\in \{\boxminus, \boxtimes, \boxplus\}, & \odot &\notin \{\boxminus, \boxtimes, \boxplus\}, & \emptyset &\notin \emptyset, \\ \{1, 2, 3\} &= \{3, 2, 1\}, \\ \{1, 1, 1\} &= \{1, 1\} = \{1\}. \end{aligned} \quad \blacktriangleleft$$

Ist die Elementzahl nicht mehr endlich, so helfen gelegentlich Pünktchen, die ein ›und so weiter‹ andeuten. So bezeichnet

$$\mathbb{N} = \{1, 2, 3, \dots\}$$

die Menge der natürlichen Zahlen⁸. Streng genommen ist ›..‹ zwar immer mehrdeutig und daher nicht exakt. Man verwendet die Pünktchen als bequeme Abkürzung aber immer dann, wenn dieses ›und so weiter‹ wirklich offensichtlich ist und nur mit viel bösem Willen falsch interpretiert werden kann ...

Schließlich werden Mengen durch Eigenschaften ihrer Elemente charakterisiert – und dies ist eigentlich auch die einzige Möglichkeit, nicht-endliche Mengen anzugeben. So bezeichnet

$$\mathbb{P} = \{n \in \mathbb{N} : n \text{ ist prim}\} = \{2, 3, 5, 7, 11, \dots\}$$

die Menge aller Primzahlen. Allgemein schreibt man

$$M = \{m : A(m)\}$$

für die Menge M aller Objekte m , die bei Einsetzen in eine *Aussageform* A ⁹ eine wahre Aussage ergeben. Will man hierbei nur die Elemente einer bestimmten Grundmenge X betrachten, so schreibt man auch kürzer¹⁰

$$\begin{aligned} M &= \{m \in X : A(m)\} \\ &:= \{m : m \in X \wedge A(m)\}. \end{aligned}$$

► **Beispiele** A. Steht $P(n)$ für die Aussageform *n ist eine Primzahl*, so ist

$$\mathbb{P} = \{n \in \mathbb{N} : P(n)\} = \{2, 3, 5, 7, 11, \dots\}.$$

B. Die Lösungsmenge der Gleichung $z^4 = 1$ im Komplexen ist

$$\{z \in \mathbb{C} : z^4 = 1\} = \{1, -1, i, -i\}. \quad \blacktriangleleft$$

Für wichtige Mengen haben sich Standardbezeichnungen eingebürgert. So bezeichnen $\mathbb{N}$, $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$, $\mathbb{C}$ die Mengen der natürlichen, ganzen, rationalen, reellen und komplexen Zahlen, die wir später kennen lernen werden.

■ Teilmengen

Eine Menge N ist *Teilmenge* einer Menge M , geschrieben

$$N \subset M,$$

⁸ Es ist reine Definitionssache, ob die natürlichen Zahlen bei 0 oder 1 beginnen. Das naive Zählen beginnt natürlicherweise bei 1.

⁹ Eine *Aussageform* enthält eine oder mehrere Variablen und geht erst durch Einsetzen konkreter Objekte in eine aristotelische Aussage über.

¹⁰ Mit $:=$ wird die linke Seite durch die rechte Seite *definiert*. Analog wird $=$ verwendet.

wenn jedes Element von N auch in M enthalten ist. Dies schließt auch die Gleichheit der beiden Mengen ein, und es gilt

$$N = M \Leftrightarrow N \subset M \wedge M \subset N.$$

Liegt eine *echte Inklusion* vor, so schreibt man ausdrücklich $N \subsetneq M$. Es gilt also

$$N \subsetneq M \Leftrightarrow N \subset M \wedge N \neq M.$$

► *Beispiele* A. Für jede beliebige Menge M gilt

$$\emptyset \subset M, \quad M \subset M.$$

B. Insbesondere ist $\emptyset \subset \emptyset$.

C. Es gilt $\mathbb{P} \subset \mathbb{N} \subset \mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R} \subset \mathbb{C}$.

D. Bezeichnet $\mathbb{P}$ die Menge der Primzahlen, so gilt ¹¹

$$\{1, 2, 3\} \not\subset \mathbb{P}, \quad \{1, 2, 3\} \subsetneq \mathbb{N}. \quad \blacktriangleleft$$

Bemerkungen a. Konsequenter wäre die Schreibweise $N \subseteq M$ für die Teilmengenbeziehung einschließlich der Gleichheit, und $N \subset M$ für die echte Teilmengenbeziehung – analog zum Gebrauch von $\leq$ und $<$ für reelle Zahlen. Die hier verwendete Notation ist aber der allgemeine Brauch. Auch Mathematiker sind nicht immer konsequent!

b. Es gilt $N \subset M$, wenn

$$x \in N \rightarrow x \in M$$

immer wahr ist. Die Definition der logischen Verknüpfung $\rightarrow$ entspricht somit der Teilmengenbeziehung $\subset$. $\rightarrow$

■ Mengenoperationen

Aus Mengen können durch unterschiedliche Operationen neue Mengen gebildet werden. Die wichtigsten sind *Vereinigung*, *Durchschnitt* und *Differenz* zweier Mengen, definiert als

$$A \cup B := \{m : m \in A \vee m \in B\},$$

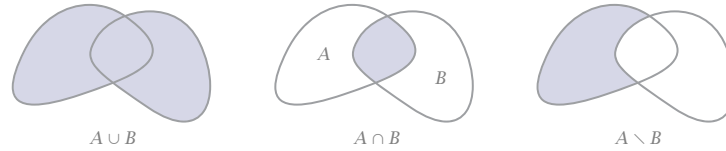
$$A \cap B := \{m : m \in A \wedge m \in B\},$$

$$A \setminus B := \{m : m \in A \wedge m \notin B\}.$$

Man nennt $A \setminus B$ auch das *relative Komplement* von B in A .

¹¹ $\not\subset$ bezeichnet die Negation von $\subset$, so wie $\notin$ die Negation von $\in$ bezeichnet.

Abb 1 Vereinigung, Durchschnitt und Differenz von Mengen



► **Beispiele** A.

$$\{H, i, l\} \cup \{e, l, f\} = \{H, i, l, f, e\},$$

$$\{H, i, l\} \cap \{e, l, f\} = \{l\},$$

$$\{H, i, l\} \setminus \{e, l, f\} = \{H, i\}.$$

B. Für eine beliebige Menge M gilt immer

$$M \cup \emptyset = M, \quad M \cap \emptyset = \emptyset,$$

$$M \setminus \emptyset = M, \quad \emptyset \setminus M = \emptyset.$$

Insbesondere gilt $\emptyset \setminus \emptyset = \emptyset$. ◀

Abbildung 1 zeigt sogenannte *Venn-Diagramme*, in denen Mengen und deren Relationen durch Bereiche in der Zeichenebene dargestellt werden. Diese sind als *Veranschaulichung* und zur *Verifizierung* sehr nützlich, ersetzen aber keinen *Beweis*, wenn es um Mengenidentitäten geht. Ein solcher Beweis wird meistens mit einer *Mengentafel* wie im Beweis des nächsten Satzes geführt.

Oft betrachtet man Teilmengen einer festen Grundmenge X . Die Differenz einer Teilmenge $M \subset X$ zur Grundmenge bezeichnet man als das *Komplement* von M bezüglich X , geschrieben

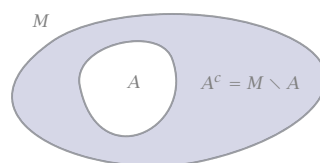
$$\complement_X M := X \setminus M = \{m : m \in X \wedge m \notin M\}.$$

Ist X aus dem Kontext bekannt und sind keine Missverständnisse zu befürchten, so schreibt man auch kürzer

$$M^c := \{m \in X : m \notin M\}.$$

Abb 2

Komplement
einer
Teilmenge



Für diese Mengenoperationen gelten eine Fülle von Rechenregeln, von denen wir die wichtigsten notieren.

10 Rechenregeln für Mengenoperationen Für beliebige Mengen A, B, C gelten die *Distributivgesetze*

$$(A \cap B) \cup C = (A \cup C) \cap (B \cup C),$$

$$(A \cup B) \cap C = (A \cap C) \cup (B \cap C),$$

das *Verschmelzungsgesetz*

$$(A \cap B) \cup A = A = (A \cup B) \cap A,$$

sowie für Teilmengen einer gemeinsamen Grundmenge X die *Regeln von de Morgan*

$$(A \cap B)^c = A^c \cup B^c, \quad (A \cup B)^c = A^c \cap B^c,$$

und das *Doppelkomplementgesetz*

$$(A^c)^c = A. \quad \times$$

»»» *Beweis* Jede dieser Identitäten beweist man mit einer *Mengentafel*. Ein beliebiges Objekt m ist in einer Menge entweder enthalten oder nicht. Schreiben wir dafür wieder 1 respektive 0, so ergibt sich zum Beispiel das erste Verschmelzungsgesetz aus folgender Mengentafel:

A	B	$A \cap B$	$(A \cap B) \cup A$	$A \cup B$	$(A \cup B) \cap A$
1	1	1	1	1	1
1	0	0	1	1	1
0	1	0	0	1	0
0	0	0	0	0	0

Analog beweist man alle übrigen Identitäten. »»»

Bemerkung Die Analogie der Rechenregeln für logische und mengentheoretischen Operationen ist natürlich kein Zufall. Mengentafeln und Wahrheitstafeln sind *äquivalent*, wenn man die Mengensymbole M und M^c als Lösungsmengen der Aussagen $m \in M$ und $m \notin M$ interpretiert und $\cap, \cup, ^c$ durch $\wedge, \vee, \neg$ ersetzt. ∞

■ Potenzmenge

Ist M eine beliebige Menge, so bezeichnet man die Menge aller ihrer *Teilmengen* als die *Potenzmenge* von M , geschrieben $\mathcal{P}(M)$. Eine Menge, deren Elemente selbst Mengen sind, wird auch als *Mengensystem* bezeichnet.

► **Beispiel** A. Für $M = \{0\}$ und $N = \{0, 1\}$ ist

$$\mathcal{P}(M) = \{\emptyset, \{0\}\} = \{\emptyset, M\},$$

$$\mathcal{P}(N) = \{\emptyset, \{0\}, \{1\}, N\}. \quad \blacktriangleleft$$

Es gilt immer

$$M \in \mathcal{P}(M), \quad \emptyset \in \mathcal{P}(M), \quad M \cap \mathcal{P}(M) = \emptyset.$$

Kein Element von M ist also ein Element von $\mathcal{P}(M)$. Vielmehr gilt

$$m \in M \Rightarrow \{m\} \in \mathcal{P}(M).$$

Die Elemente von M sind als Ein-Element-Mengen *verpackt* in $\mathcal{P}(M)$ enthalten.

► **Noch ein Beispiel**

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

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

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

Stellt man sich die leere Menge $\emptyset$ als Sack vor, der nichts enthält, so ist $\{\emptyset\}$ ein Sack, der einen leeren Sack enthält, und $\{\{\emptyset\}\}$ ein Sack, der einen Sack enthält, der einen leeren Sack enthält. Dies sind drei sehr unterschiedliche Dinge! $\blacktriangleleft$

■ Kartesisches Produkt

Sind A und B zwei beliebige Mengen, so ist deren *kartesisches Produkt* die Menge aller *geordneten Paare*, die sich aus Elementen von A an erster Position und Elementen von B an zweiter Position bilden lassen. In Symbolen:

$$A \times B := \{(a, b) : a \in A, b \in B\}.$$

Hierbei ist ein *geordnetes Paar* (a, b) eindeutig durch seine beiden *Komponenten* a und b und deren Reihenfolge festgelegt ist¹². Ein Paar ist also etwas wesentlich anderes als eine Menge mit zwei Elementen. Für $a \neq b$ ist

$$\{a, b\} = \{b, a\}, \quad (a, b) \neq (b, a).$$

Ist eine der beteiligten Mengen leer, so ist das kartesische Produkt vereinbarungsgemäß ebenfalls leer. Für jede Menge M gilt also

$$M \times \emptyset = \emptyset \times M = \emptyset.$$

¹² Wir begnügen uns mit einer naiven Definition eines geordneten Paares. Denn was ist genau eine erste und zweite Position? Eine mengentheoretische Definition wäre übrigens $(a, b) := \{a, \{a, b\}\}$.

► *Beispiele* A. Mit $A = \{\ominus, \boxtimes\}$ und $B = \{\ominus, \boxplus\}$ ist

$$A \times B = \{(\ominus, \ominus), (\ominus, \boxplus), (\boxtimes, \ominus), (\boxtimes, \boxplus)\},$$

$$B \times A = \{(\ominus, \ominus), (\ominus, \boxtimes), (\boxplus, \ominus), (\boxplus, \boxtimes)\}.$$

Insbesondere ist $A \times B \neq B \times A$.

B. Für eine nichtleere Menge A ist

$$A^2 := A \times A = \{(a, b) : a, b \in A\}. \quad \blacktriangleleft$$

Graphisch kann man das kartesische Produkt $A \times B$ durch ein Rechteck veranschaulichen, dessen Seiten die Mengen A und B repräsentieren. Jeder Punkt des Rechtecks stellt dann ein Koordinatenpaar (a, b) mit $a \in A$ und $b \in B$ dar.

Analog werden Produkte aus mehr als zwei Mengen definiert. So ist

$$A^3 := A \times A \times A = \{(a, b, c) : a, b, c \in A\}$$

die Menge aller *Tripel* (a, b, c) mit Komponenten in A . Allgemein ist, für jede natürliche Zahl $n \geq 1$,

$$A^n := \{(a_1, a_2, \dots, a_n) : a_1, \dots, a_n \in A\}$$

die Menge aller n -Tupel mit Komponenten in A .

► *Beispiele* A. Der dreidimensionale Raum unserer Anschauung wird beschrieben durch die Menge aller geordneten Tripel reeller Zahlen,

$$\mathbb{R}^3 = \mathbb{R} \times \mathbb{R} \times \mathbb{R} = \{(x, y, z) : x, y, z \in \mathbb{R}\}.$$

B. Der n -dimensionale reelle Vektorraum $\mathbb{R}^n$ besteht aus allen n -Tupeln reeller Zahlen,

$$\mathbb{R}^n = \{(x_1, \dots, x_n) : x_1, \dots, x_n \in \mathbb{R}\}. \quad \blacktriangleleft$$

1.3

Relationen

Elemente einer Menge M können auf vielfältige Weise miteinander in Beziehung gebracht werden. Beispielsweise können zwei Teilmengen einer festen Menge mindestens ein Element gemeinsam haben, oder zwei natürliche Zahlen können einen gemeinsamen Teiler haben. Um solche Beziehungen aber ohne Bezug auf eine semantische Interpretation zu beschreiben, betrachtet man *jede* Teilmenge des kartesischen Produktes $M \times M$ als eine *Relation* in der Menge M .

Definition Eine *Relation* auf einer Menge M ist eine Teilmenge $R \subset M \times M$. Man schreibt

$$aRb :\Leftrightarrow (a, b) \in R$$

und sagt, *a* und *b* stehen in Relation *R* zueinander. ✕

Wir benötigen folgende Eigenschaften von Relationen.

Definition Eine Relation $R \subset M \times M$ heißt *total*, falls zwei beliebige Elementen immer in Relation *R* stehen, also $aRb \vee bRa$ gilt. Sie heißt

- *reflexiv*, falls aRa ,
 - *symmetrisch*, falls $aRb \Rightarrow bRa$,
 - *antisymmetrisch*, falls $aRb \wedge bRa \Rightarrow a = b$,
 - *transitive*, falls $aRb \wedge bRc \Rightarrow aRc$,
- für alle $a, b, c \in M$ gilt. ✕

Beispiele geben wir gleich. — Besonders wichtig sind für uns *Äquivalenz-* und *Ordnungsrelationen*.

■ Äquivalenzrelationen

Definition Eine *Äquivalenzrelation* ist eine reflexive, symmetrische und transitive Relation und wird im Allgemeinen mit $\sim$ bezeichnet. ✕

Für eine Äquivalenzrelation $\sim$ auf M gilt also

$$a \sim a, \quad a \sim b \Rightarrow b \sim a, \quad a \sim b \wedge b \sim c \Rightarrow a \sim c.$$

Gilt $a \sim b$, so heißen *a* und *b* *äquivalent*, und die Menge

$$[a] := \{b \in M : b \sim a\}$$

aller zu *a* äquivalenten Elemente heißt die *Äquivalenzklasse* von *a*.

- 11 ➤ **Beispiele** A. Die Gleichheitsrelation $=$ ist auf jeder Menge eine Äquivalenzrelation. Sie ist die feinstmögliche, denn ihre Äquivalenzklassen sind sämtlich Ein-Punkt-Mengen:

$$[a] = \{a\}, \quad a \in M.$$

- B. Sei $p \in \mathbb{N}$ mit $p \geq 2$. Dann definiert

$$n \equiv m \bmod p \Leftrightarrow p \mid n - m \Leftrightarrow p \text{ teilt } n - m$$

eine Äquivalenzrelation auf der Menge $\mathbb{Z}$ der ganzen Zahlen _{A-24}. Die Äquivalenzklassen dieser Relation sind die Restklassen bezüglich Division durch p , von denen es p gibt. Anders ausgedrückt,

$$[n] = \{n + kp : k \in \mathbb{Z}\}.$$

- C. Auf $\mathbb{N} \times \mathbb{N}$ definiert

$$(n, m) \sim (p, q) \Leftrightarrow nq = mp$$

eine Äquivalenzrelation. Die Äquivalenzklasse $[(n, m)]$ besteht aus allen Repräsentanten derselben rationalen Zahl n/m , wenn man diese als Zahlenpaar schreibt.

- D. In der euklidischen Ebene $\mathbb{R} \times \mathbb{R}$ definiert

$$(x_1, y_1) \sim (x_2, y_2) \Leftrightarrow x_1 = x_2$$

eine Äquivalenzrelation. Ihre Äquivalenzklassen sind die Punkte auf senkrechten Geraden. ◀

- 12 **Satz** Ist $\sim$ eine Äquivalenzrelation auf der Menge M , so bildet ¹³

$$M/\sim := \{[a] : a \in M\},$$

genannt die *Restklassenmenge von M modulo $\sim$* , eine *Zerlegung* von M : Die Vereinigung aller Äquivalenzklassen ist M , und zwei Äquivalenzklassen sind entweder gleich oder disjunkt. Es gilt also

$$M = \bigcup_{a \in M} [a] \quad \text{und} \quad [a] \cap [b] \neq \emptyset \Rightarrow [a] = [b]. \quad \times$$

Mit anderen Worten, jedes Element von M gehört zu genau einer Äquivalenzklasse in $M/\sim$.

¹³ Erinnerung: Eine Menge enthält ein Element nur *einmal*, auch wenn es mehrmals notiert wird.

⟨⟨⟨⟨ Aufgrund der Reflexivität einer Äquivalenzrelation ist $a \in [a]$ für alle $a \in M$ und deshalb

$$M \subset \bigcup_{a \in M} [a] \subset M.$$

Also sind diese beiden Mengen gleich.

Angenommen, es ist nun $[a] \cap [b] \neq \emptyset$. Somit gibt es ein $c \in [a] \cap [b]$, und es gilt $c \sim a$ und $c \sim b$. Also ist auch $a \sim b$ aufgrund der Symmetrie und Transitivität einer Äquivalenzrelation. Daraus aber folgt, dass $[a] = [b]$. ⟩⟩⟩⟩

■ Ordnungsrelationen

Definition Eine *Ordnungsrelation* ist eine reflexive, antisymmetrische und transitive Relation und wird mit $\leq$ bezeichnet, falls keine speziellere Bezeichnung vereinbart wird. Ist die Relation außerdem total, so spricht man von einer *totalen Ordnungsrelation*. ✕

Für eine Ordnungsrelation $\leq$ auf M gilt also

$$a \leq a, \quad a \leq b \wedge b \leq a \Rightarrow b = a, \quad a \leq b \wedge b \leq c \Rightarrow a \leq c$$

für alle $a, b, c \in M$. Für eine totale Ordnung gilt außerdem $a \leq b \vee b \leq a$ für beliebige $a, b \in M$. Ist die Ordnung nicht total, so gibt es mindestens zwei Elemente, die nicht in Relation stehen.

► **Beispiele** A. Die Gleichheitsrelation ist eine Ordnungsrelation. Sie ist aber nicht total, wenn M mehr als ein Element enthält.

B. Die übliche $\leq$ -Relation auf $\mathbb{N}$ ist eine totale Ordnungsrelation.

C. Dasselbe gilt für $\leq$ auf $\mathbb{Z}$, $\mathbb{Q}$ und $\mathbb{R}$.

D. Auf der Potenzmenge einer beliebigen Menge M definiert die Teilmengenbeziehung $\subset$ eine Ordnung. Diese ist allerdings *nicht total*, wenn M mehr als ein Element enthält.

E. Auf $\mathbb{N}$ definiert die *Teiler-Relation* $n \mid m :\Leftrightarrow n$ teilt m eine nicht-totale Ordnung. ◀

Für eine Ordnungsrelation $\leq$ erklärt man außerdem

$$a < b :\Leftrightarrow a \leq b \wedge a \neq b$$

sowie $a \geq b :\Leftrightarrow b \leq a$ und $a > b :\Leftrightarrow b < a$. Dann gilt der folgende

13 Trichotomiesatz Ist $\leq$ eine totale Ordnung auf M , so gilt für je zwei Elemente $a, b \in M$ immer eine und nur eine der drei Relationen

$$a < b, \quad a = b, \quad a > b.$$

Ist umgekehrt $<$ eine transitive Relation auf M , so dass für je zwei beliebige Elemente diese Trichotomie gilt, so definiert

$$a \leq b \quad :\Leftrightarrow \quad a < b \vee a = b$$

eine totale Ordnung auf M . $\times$

⟦⟦⟦ Beweis der ersten Aussage: Gilt zum Beispiel $a < b$, so ist *per definitionem* $a \neq b$. Auch kann nicht $a > b$ gelten, denn dann gilt auch $a \leq b \wedge a \geq b$ und damit wieder $a = b$. Gilt dagegen $a = b$, so kann ebenfalls *per definitionem* weder $a < b$ noch $a > b$ gelten. Somit gilt immer *genau eine* der drei Aussagen.

Beweis der zweiten Aussage: Wegen $a = a$ gilt auch $a \leq a$. Also ist $\leq$ reflexiv. Gilt $a \leq b \wedge b \leq a$, so gilt

$$\begin{aligned} & (a < b \vee a = b) \wedge (b < a \vee b = a) \\ \Leftrightarrow & (a < b \wedge b < a) \vee (a = b) \\ \Leftrightarrow & (a = b), \end{aligned}$$

da $(a < b \wedge b < a)$ aufgrund der Trichotomie falsch ist. Also ist $\leq$ antisymmetrisch. Und da $<$ transitiv ist, ist auch $\leq$ transitiv. Schließlich folgt aus der Trichotomie

$$\begin{aligned} & (a < b) \vee (a = b) \vee (a > b) \\ \Leftrightarrow & (a < b \vee a = b) \vee (a = b \vee a > b) \\ \Leftrightarrow & (a \leq b) \vee (a \geq b). \end{aligned}$$

Also ist $\leq$ total. ⟦⟦⟦

► **Beispiel** Ist $\leq$ eine totale Ordnung auf M , so wird auf $M \times M$ durch

$$(a, b) \leq (c, d) \quad :\Leftrightarrow \quad (a < c) \vee (a = c \wedge b \leq d)$$

eine totale Ordnung definiert, genannt *lexikographische Ordnung* A-23. ◀

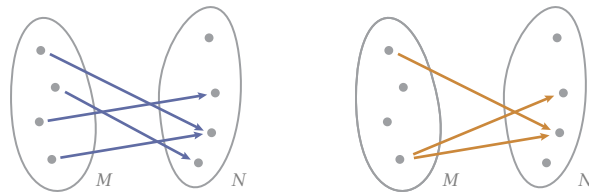
1.4

Abbildungen

Mengen kann man aufeinander abbilden. Der Begriff der *Abbildung* ist von ebenso fundamentaler und weitreichender Bedeutung wie der Begriff der Menge. Auch hier werden wir uns mit einer naiven Definition begnügen.

- 14 **Naive Definition** Seien M und N zwei beliebige Mengen. Eine *Abbildung* oder *Funktion* f von M nach N ist eine Vorschrift, die jedem Element $a \in M$

Abb 3

Funktion und
Nicht-Funktion

genau ein Element $b \in N$ zuordnet. In Symbolen:

$$f: M \rightarrow N, \quad a \mapsto b = f(a). \quad \times$$

Diese Definition ist ›naiv‹, da wir an eine anschauliche Bedeutung der Begriffe ›Vorschrift‹ und ›zuordnen‹ appellieren, ohne diese zu präzisieren. Sie ist sogar eher irreführend, da sie suggeriert, dass jede Abbildung in Gestalt einer Formel daherkommt. Dem ist aber nicht so! Sehr viele wichtige mathematische Funktionen sind nur *implizit* definiert, ohne einen expliziten formelmäßigen Ausdruck. Man kann sogar *beweisen*, dass es für sie solche Formeln nicht gibt. Dies gilt zum Beispiel für die Lösungen der meisten Differenzialgleichungen.

- 15 ➤ **Beispiele** A. Die aus der Schule vertraute lineare Funktion mit Steigung $m \in \mathbb{R}$ und Ordinatenabschnitt $b \in \mathbb{R}$ ist

$$\lambda: \mathbb{R} \rightarrow \mathbb{R}, \quad t \mapsto \lambda(t) = mt + b.$$

- B. Die klassische quadratische Funktion ist

$$p: \mathbb{R} \rightarrow \mathbb{R}, \quad t \mapsto p(t) = t^2.$$

- C. Die Funktion

$$\varphi: \mathbb{N} \rightarrow \mathbb{N} \cup \{0\}, \quad n \mapsto \varphi(n) = \text{card} \{p \in \mathbb{P} : p \leq n\}$$

liefert für jede natürliche Zahl n die Anzahl der Primzahlen $\leq n$. Diese kann nicht durch eine einfache Formel angegeben werden. ◀

Im Falle endlicher Mengen M und N ist eine Funktion $f: M \rightarrow N$ im Prinzip immer durch ein Pfeildiagramm wie in Abbildung 3 angebbbar. Die Punkte bezeichnen die Elemente der Mengen M und N , und die Pfeile deren Zuordnung durch die Abbildung f . Eine Funktion liegt genau dann vor, wenn von *jedem* Punkt der Menge M *genau ein* Pfeil ausgeht. Andererseits dürfen Punkte in N von keinem, einem, oder mehreren Pfeilen getroffen werden. Es dürfen sogar *alle* Punkte aus M auf denselben Punkt in N abgebildet werden.

Nun etwas Terminologie. Mit der etwas schwerfällig erscheinenden Notation

$$f: M \rightarrow N,$$

$$a \mapsto b = f(a)$$

werden alle Bestandteile notiert, die eine Funktion ausmachen: f ist der *Name* der Funktion, M ihr *Definitionsbereich* und N ihr *Wertebereich*. Jeden Punkt $a \in M$ bildet f auf den *Funktionswert* $b = f(a) \in N$ ab, den man auch den *Bildpunkt* von a unter f nennt. Der Punkt a selbst heißt *Urbild* von b unter f . Einem Punkt $a \in M$ ist immer *genau ein Bildpunkt* $b \in N$ zugeordnet. Ein Punkt $b \in N$ kann aber *mehrere Urbilder* haben – siehe Abbildung 3.

Es ist üblich, den einfachen Pfeil $\rightarrow$ zwischen Definitions- und Wertebereich zu setzen, und den abgeschlossenen Pfeil $\mapsto$ zwischen deren Elementen, um diese verschiedenen Ebenen zu unterscheiden.

Natürlich verwendet man auch kürzere Schreibweisen. Wird der Name einer Funktion nicht benötigt, so schreibt man auch kürzer

$$\mathbb{R} \rightarrow \mathbb{R}, \quad t \mapsto t^2.$$

Ist auch der Definitionsbereich im gegebenen Kontext unwesentlich oder bekannt, so schreibt man schlicht $t \mapsto t^2$. Ist andererseits die ›Abbildungsvorschrift‹ nicht von Belang, so schreibt man $f: M \rightarrow N$ für irgendeine Abbildung von M nach N .

Das *Bild von M* unter einer Abbildung $f: M \rightarrow N$ ist die Menge

$$f(M) := \{f(a) : a \in M\} \subset N$$

aller Bildpunkte unter f . Dies ist etwas wesentlich anderes als der Wertebereich. Während dieser meist leicht festgelegt werden kann, ist das genaue Bild einer Menge unter einer Funktion gelegentlich schwer zu bestimmen.

16 ➤ *Beispiel* Mit den Bezeichnungen des vorangehenden Beispiels $_{15}$ ist

$$\lambda(\mathbb{R}) = \begin{cases} \{b\}, & m = 0, \\ \mathbb{R}, & m \neq 0, \end{cases}$$

und

$$p(\mathbb{R}) = \{x \in \mathbb{R} : x \geq 0\}.$$

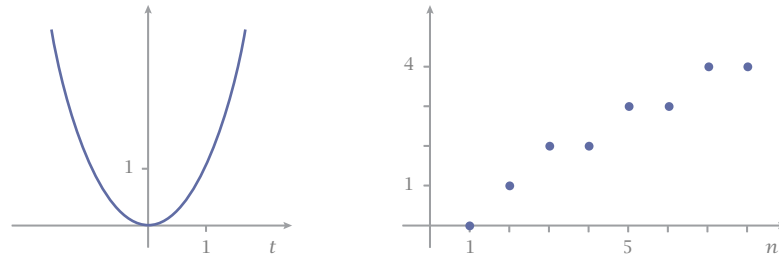
Da es unendlich viele Primzahlen gibt, ist außerdem

$$\varphi(\mathbb{N}) = \{0, 1, 2, \dots\} = \mathbb{N} \cup \{0\}. \quad \blacktriangleleft$$

■ Graph

Jede Funktion $f: M \rightarrow N$ bestimmt eindeutig ihren *Graph*

$$\Gamma(f) := \{(a, b) \in M \times N : a \in M \wedge b = f(a) \in N\}.$$

Abb 4 Graphen der Funktionen p und φ aus Beispiel 15

Dies ist also eine *Relation* in $M \times N$, bestehend aus allen Paaren der Form $(a, f(a))$ mit $a \in M$. In dieser Relation kommt also jedes Element $a \in M$ *genau einmal* als erste Komponente vor.

Handelt es sich bei M und N um die reelle Gerade, so können wir Graphen oft zeichnerisch in gewohnter Weise wie in Abbildung 4 darstellen.

Umgekehrt bestimmt ein Graph $\Gamma(f)$ einer Funktion eindeutig die zugrunde liegende Funktion f . Denn zu jedem $a \in M$ existiert genau ein $b \in N$, so dass $(a, b) \in \Gamma(f)$, und dieses b ist der eindeutige Bildpunkt von a unter f . Aufgrund dieses eindeutigen Zusammenhangs kann man den Begriff der Funktion auf den Begriff der Relation zurückführen.

Mengentheoretische Definition Eine *Abbildung* $f: M \rightarrow N$ ist eine Relation in $M \times N$, in der es zu jedem $a \in M$ genau ein $b \in N$ gibt, so dass $a f b$ gilt. Eine solche Relation wird auch *funktionale Relation* genannt. $\times$

Diese Definition vermeidet Begriffe wie ›Vorschrift‹ und ›zuordnen‹ und damit die Assoziation von ›Funktion‹ mit ›Formel‹, ist dafür aber weniger anschaulich.

In jedem Fall ist eine Funktion ein Objekt, das aus einem *Definitionsbereich*, einem *Wertebereich* und einer *funktionalen Relation* zwischen beiden besteht. Diese drei Bestandteile gehören genannt, wenn sie nicht aus dem Kontext bekannt sind oder für die jeweilige Betrachtung keine Rolle spielen. Zwei Funktionen sind *gleich* dann nur dann, wenn sie in *allen* diesen Aspekten gleich sind. Das heißt,

$$f: M \rightarrow N, \quad g: U \rightarrow V$$

sind *gleich* dann und nur dann, wenn $M = U$ und $N = V$ und

$$f(a) = g(a) \quad \text{für alle } a \in M = U.$$

Die Wahl des Wertebereiches wird allerdings erst dann bedeutsam, wenn dieser mit einer zusätzlichen Struktur wie zum Beispiel einer Topologie versehen ist.

■ Standardfunktionen

Zunächst bemerken wir, dass es *keine Abbildung* einer nichtleeren Menge M in die leere Menge gibt. Denn einem Punkt $a \in M$ kann kein Bildpunkt zugeordnet werden. Es gibt aber umgekehrt immer genau eine Abbildung der leeren Menge in eine beliebige nichtleere Menge N , die *leere Abbildung*

$$e : \emptyset \rightarrow N.$$

Diese wird uns allerdings nicht sehr beschäftigen. — Nun einige wichtige Standardfunktionen.

a. Die Abbildung

$$id_M : M \rightarrow M, \quad a \mapsto a$$

heißt die *Identität* auf M . Ist der Definitionsbereich klar, schreibt man auch id .

b. Ist $A \subset M$, so heißt

$$i : A \rightarrow M, \quad a \mapsto a$$

die *Inklusionsabbildung* oder *Einbettung* von A in M . Es ist $i = id_M$ genau dann, wenn $A = M$.

c. Ist $A \subset M$, so heißt

$$\chi_A : M \rightarrow \{0, 1\}, \quad \chi_A(a) = \begin{cases} 1 & \text{für } a \in A, \\ 0 & \text{für } a \in A^c, \end{cases}$$

die *charakteristische Funktion* oder *Indikatorfunktion* von A auf M .

d. Ist eine Funktion $f : M \rightarrow N$ gegeben und $A \subset M$, so heißt

$$f|_A : A \rightarrow N, \quad a \mapsto f(a)$$

die *Einschränkung* von f auf A . Es ist $f|_A = f$ genau dann, wenn $A = M$.

■ Tupel, Folgen, Operationen

Wir erwähnen noch einige Funktionen spezieller Art, die unter anderen Namen bekannt sind. Hierbei sei X eine beliebige nichtleere Menge.

a. Für jede natürliche Zahl $n \geq 1$ sei $\mathbb{A}_n := \{1, 2, \dots, n\}$ die Menge der ersten n natürlichen Zahlen. Eine Funktion

$$f : \mathbb{A}_n \rightarrow X$$

wird vollständig beschrieben durch ihre n Funktionswerte

$$x_k := f(k), \quad 1 \leq k \leq n.$$

Diese werden bequemer in Form eines *n -Tupels* $(x_1, \dots, x_n)$ angegeben. Der Funktionswert $f(k)$ ist dann die k -te *Koordinate* oder *Komponente* x_k des Tupels. Insbesondere ist ein 2-Tupel ein geordnetes Paar.¹⁴

n -Tupel sind also eine Kurznotation für Funktionen auf $\mathbb{A}_n$. Dementsprechend sind zwei solche Tupel gleich, wenn alle entsprechenden Komponenten gleich sind:

$$(x_1, \dots, x_n) = (y_1, \dots, y_n) \Leftrightarrow x_k = y_k, \quad 1 \leq k \leq n.$$

Die Menge aller n -Tupel mit Komponenten in X ist das *n -fache kartesische Produkt* der Menge X ,

$$X^n = \{(x_1, \dots, x_n) : x_1, \dots, x_n \in X\}.$$

Am häufigsten wird uns der *n -dimensionale reelle Raum*

$$\mathbb{R}^n = \{(x_1, \dots, x_n) : x_1, \dots, x_n \in \mathbb{R}\}.$$

begegnen.

b. Entsprechend wird eine Funktion $f: \mathbb{N} \rightarrow X$ vollständig beschrieben durch ihre Funktionswerte

$$x_k := f(k), \quad k = 1, 2, \dots$$

Man spricht von einer *Folge* in X , und schreibt sie in der Form

$$(x_1, x_2, \dots) = (x_k)_{k \geq 1} = (x_k).$$

Folgen werden wir genauer in Kapitel 4 studieren.

c. Eine Funktion

$$\star : X \times X \rightarrow X$$

ordnet jedem Paar $(a, b) \in X \times X$ ein neues Element $c = \star(a, b) \in X$ zu. Dies kann man als eine *binäre Operation* oder *innere Verknüpfung* auf X auffassen. Die übliche Schreibweise hierfür ist

$$c = a \star b.$$

Eine binäre Operation $\star$ auf X heißt *kommutativ*, falls

$$a \star b = b \star a$$

¹⁴ Dies ist allerdings kein Ersatz für unsere naive Definition des Paares! Denn der mengentheoretische Begriff der Funktion basiert auf dem Begriff des kartesischen Produktes, und dieses wiederum haben wir mit dem Begriff des geordneten Paares erklärt.

für alle $a, b \in X$. Sie heißt *assoziativ*, falls

$$a \star (b \star c) = (a \star b) \star c$$

für alle $a, b, c \in X$.

► *Beispiel* Zwei aus früher Kindheit vertraute binäre Operationen sind die Addition und die Multiplikation natürlicher Zahlen. Beide sind assoziativ und kommutativ. ◀

■ Komposition

Zwei Abbildungen können verknüpft oder hintereinander ausgeführt werden, wenn der Wertebereich des ersten Abbildung im Definitionsbereich der zweiten Abbildung enthalten ist. Wir betrachten also zwei Abbildungen $f: M \rightarrow N$ und $g: N \rightarrow O$, wofür man auch anschaulicher

$$M \xrightarrow{f} N \xrightarrow{g} O$$

schreibt. Dann ist die *Komposition* von f mit g definiert als die Abbildung

$$g \circ f: M \rightarrow O, \quad a \mapsto (g \circ f)(a) := g(f(a)).$$

Gelesen wird dies *g nach f* oder *g kringel f*. Die Komposition operiert *immer von rechts nach links*: zuerst wird f angewandt, danach g , so wie ja auch f auf das rechts stehende Argument a angewandt wird.

Die Komposition ist *immer assoziativ*. Ist also die Verknüpfung dreier Abbildungen überhaupt definiert, so ist immer

$$(f \circ g) \circ h = f \circ (g \circ h),$$

weshalb man die Klammern auch ganz weglassen kann. Dagegen ist $f \circ g$ im Allgemeinen etwas anderes als $g \circ f$, wie man sich leicht anhand eines Beispiels überlegt. Die Komposition ist also *nicht kommutativ*.

► *Beispiel* Bezeichnet $F(M)$ den Raum aller Abbildungen von $M \rightarrow M$, so definiert die Komposition

$$\circ: F(M) \times F(M) \rightarrow F(M), \quad (g, f) \mapsto g \circ f$$

eine *Operation* auf $F(M)$. ◀

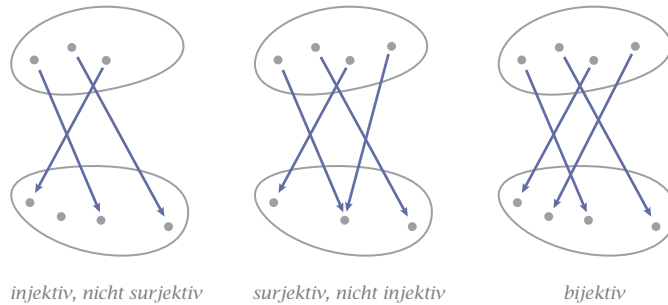
■ Injektion, Surjektion, Bijektion

Es folgen die drei wichtigsten mengentheoretischen Eigenschaften, die eine Abbildung aufweisen kann.

Definition Eine Abbildung $f: M \rightarrow N$ heißt

Abb 5

Injektion,
Surjektion,
Bijektion



injektiv, nicht surjektiv

surjektiv, nicht injektiv

bijektiv

- **injektiv**, wenn jeder Punkt in N höchstens ein Urbild besitzt,
- **surjektiv**, wenn jeder Punkt in N mindestens ein Urbild besitzt,
- **bijektiv**, wenn sie injektiv und surjektiv ist. ✕

Eine Abbildung $f: M \rightarrow N$ ist somit surjektiv, wenn $f(M) = N$. Man sagt dann auch, f bildet M auf N ab. Sie ist injektiv, wenn keine zwei Urbilder in M dasselbe Bild haben. Mit anderen Worten, für alle $a, b \in M$ gilt

$$f(a) = f(b) \Rightarrow a = b.$$

Sie ist bijektiv, wenn jeder Punkt in N genau ein Urbild besitzt.

- 17 **Satz** Eine Abbildung $f: M \rightarrow N$ ist bijektiv genau dann, wenn es eine Abbildung $\varphi: N \rightarrow M$ gibt, so dass

$$\varphi \circ f = id_M, \quad f \circ \varphi = id_N.$$

In diesem Fall ist φ eindeutig bestimmt. ✕

⟨⟨⟨ Beweis $\Rightarrow$ Ist f bijektiv, so gibt es zu jedem $b \in N$ genau ein $a \in M$ mit $f(a) = b$. Wir können dadurch eine Abbildung $\varphi: N \rightarrow M$ definieren, welche die gewünschten Eigenschaften besitzt.

$\Leftarrow$ Für jedes $b \in N$ gilt wegen $f \circ \varphi = id_N$

$$b = id_N(b) = (f \circ \varphi)(b) = f(\varphi(b)).$$

Also ist f surjektiv. Gilt andererseits $f(a_1) = f(a_2)$ für zwei Elemente in M , so folgt mit $\varphi \circ f = id_M$

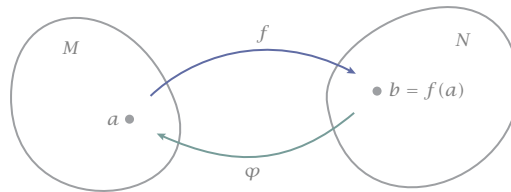
$$a_1 = \varphi(f(a_1)) = \varphi(f(a_2)) = a_2.$$

Also ist f auch injektiv. Somit ist f bijektiv.

Bleibt noch die Eindeutigkeit von φ zu zeigen. Ist $\psi: N \rightarrow M$ eine weitere Abbildung mit den Eigenschaften $\psi \circ f = id_M$ und $f \circ \psi = id_N$, so folgt

Abb 6

Abbildung
und Umkehr-
abbildung



$$\begin{aligned}\psi &= \psi \circ id_N = \psi \circ (f \circ \varphi) \\ &= (\psi \circ f) \circ \varphi \\ &= id_M \circ \varphi \\ &= \varphi.\end{aligned}$$

Also ist $\psi = \varphi$, und es gibt nur eine solche Abbildung. $\gggg$

■ Umkehrabbildung

Eine bijektive Abbildung $f: M \rightarrow N$ ist also umkehrbar₁₇, und wir können ihre **Umkehrabbildung** f^{-1} definieren als die eindeutig bestimmte Abbildung $f^{-1}: N \rightarrow M$ mit der Eigenschaft, dass

$$f^{-1} \circ f = id_M, \quad f \circ f^{-1} = id_N.$$

Das folgende Lemma ist als Aufgabe überlassen_{A-28}.

18 Lemma Die Komposition umkehrbarer Abbildungen ist umkehrbar, und es gilt

$$(g \circ f)^{-1} = f^{-1} \circ g^{-1}. \quad \times$$

19 ➤ **Beispiele** A. Die Abbildung

$$p: \mathbb{R} \rightarrow \mathbb{R}, \quad t \mapsto p(t) = t^2$$

ist nicht umkehrbar, da sie nicht injektiv ist. So ist $p(-1) = p(1) = 1$.

B. Die Abbildung

$$q: [0, \infty) \rightarrow \mathbb{R}, \quad t \mapsto q(t) = t^2$$

ist bijektiv auf ihr Bild $[0, \infty)$ und damit umkehrbar. Ihre Umkehrfunktion ist die Quadratwurzelfunktion.

C. Die Umkehrfunktion der Exponentialfunktion,

$$\exp: \mathbb{R} \rightarrow (0, \infty), \quad t \mapsto e^t,$$

ist die Logarithmusfunktion ?? $\lll$