

1 Elementare Logik

1. Aussagenlogik

Unter einer **Aussage** verstehen wir einen grammatikalisch korrekten Satz, dem ein Wahrheitswert zugewiesen werden kann. Als Wahrheitswerte sind dabei ausschließlich „wahr“ und „falsch“ zugelassen.

Als typische Bezeichnungen für Aussagen verwenden wir Buchstaben $A, B, C, \dots$, für die Wahrheitswerte die Buchstaben

w für „wahr“ und f für „falsch“.

Beispiel. Die folgenden Sätze sind korrekt gebildet:

A : 9 ist eine Primzahl

B : 1729 ist die Summe zweier Kubikzahlen

C : 1234567891 ist eine Primzahl

D : Jede gerade Zahl, die größer als Zwei ist, ist die Summe zweier Primzahlen

Aber sind sie auch Aussagen? Offenbar ist A falsch. B lässt sich leicht überprüfen: es ist $1729 = 1^3 + 12^3 = 9^3 + 10^3$. Es ist klar, dass man auch C mit etwas mehr Mühe ebenfalls nachprüfen kann (tatsächlich ist C wahr), aber was ist mit D ? Die Antwort kennt kein Mensch, aber dennoch ist die Aussage entweder wahr oder falsch. Ein Satz kann durchaus einen Wahrheitswert haben, auch wenn wir ihn vielleicht nicht kennen!

Andererseits sind zum Beispiel Fragesätze keine Aussagen. Schwieriger wird es mit Sätzen der Art „Dieser Satz ist falsch“, denn welchen Wahrheitswert wir ihm auch zuweisen wollen, widersprechen wir uns selbst. Ein solches Paradoxon ist also nach unserer Auffassung keine Aussage und wir gar nicht erst zum Aussagenkalkül zugelassen. Wir werden uns auch im wesentlichen auf mathematische Aussagen beschränken, um Schwierigkeiten mit mehrfachen Bedeutungen von Wörtern aus dem Weg zu gehen.

Aus gegebenen Aussagen kann man durch geeignete logische Operationen neue Aussagen bilden. Dabei gibt man mit Hilfe einer **Wahrheitstafel** an, welchen Wahrheitswert die neue Aussage haben soll.

Die **Negation** einer Aussage A ist die Aussage „nicht A “, bezeichnet durch die symbolische Schreibweise „ $\neg A$ “, mit der Wahrheitstafel

A	$\neg A$
w	f
f	w

Hier stehen in der ersten Spalte die möglichen Wahrheitswerte für A und in der zweiten die dazugehörigen Werte der Aussage $\neg A$.

Die **Disjunktion** der Aussagen A, B ist die Aussage „ A oder B “, symbolisch abgekürzt durch „ $A \vee B$ “. Sie wird beschrieben durch die Wahrheitstafel

A	B	$A \vee B$		$A \backslash B$	w	f
w	w	w		w	w	w
w	f	w	oder auch	f	w	f
f	w	w		w	w	w
f	f	f		f	w	f

(In der zweiten Tafel steht in der Position (a, b) des quadratischen Felds der Wahrheitswert der Aussage $A \vee B$, wenn A den Wahrheitswert a und B den Wert b hat).

Die Aussage $A \vee B$ ist also genau dann wahr, wenn wenigstens eine der beiden Aussagen A, B wahr ist. Dies ist das „mathematische oder“ und nicht das ausschließende „entweder – oder“ der Umgangssprache. Dieses hätte eine andere Wahrheitstafel, nämlich

A	B	entweder A oder B		$A \backslash B$	w	f
w	w	f		w	f	w
w	f	w	bzw.	f	w	f
f	w	w		w	w	w
f	f	f		f	w	f

Die **Konjunktion** zweier Aussagen A, B ist die Aussage „ A und B “, abgekürzt geschrieben als „ $A \wedge B$ “; sie ist gegeben durch die Wahrheitstafel

A	B	$A \wedge B$		$A \backslash B$	w	f
w	w	w		w	w	f
w	f	f	bzw.	f	w	f
f	w	f		w	f	f
f	f	f		f	f	f

Besonders wichtig ist die **Implikation**, also die Aussage „aus A folgt B “ oder „ A impliziert B “ oder „wenn A dann B “, in Kurzschreibweise $A \Rightarrow B$, mit zugehöriger Wahrheitstafel

A	B	$A \Rightarrow B$		$A \backslash B$	w	f
w	w	w		w	w	f
w	f	f	bzw.	f	w	w
f	w	w		w	f	f
f	f	w		f	f	f

Man nennt dann die Aussage A die *Voraussetzung* oder *Prämisse* und B die *Behauptung* oder *Conclusio*. Die Aussage „ $A \Rightarrow B$ “ ist also genau dann falsch, wenn A falsch und B wahr ist. Insbesondere kann man aus einer falschen Voraussetzung alles schließen!

Die **Äquivalenz** zweier Aussagen, geschrieben „ $A \Leftrightarrow B$ “ und gesprochen „genau dann A , wenn B “, ist erklärt durch

$$(A \Rightarrow B) \wedge (B \Rightarrow A),$$

hat also die Wahrheitstafel

A	B	$A \Leftrightarrow B$		$A \backslash B$	w	f
w	w	w		w	w	f
w	f	f	bzw.	f	w	f
f	w	f		w	f	w
f	f	w		f	f	w

Disjunktion, Konjunktion, Implikation und Äquivalenz verknüpfen Aussagen, daher nennt man sie auch **Junktoren**.

Verwendet man mehrere Operationen hintereinander, so muss man Klammern setzen: $A \vee (B \wedge C)$ ist nicht das gleiche wie $(A \vee B) \wedge C$. Dabei gilt allerdings die Konvention, dass das Negationssymbol am stärksten bindet, so dass etwa $\neg A \vee B$ zu lesen ist als $(\neg A) \vee B$.

Satz 1.1. *Sei A eine Aussage. Dann gelten die folgenden Regeln:*

1. $\neg(\neg A) \Leftrightarrow A$ (Regel von der doppelten Verneinung)
2. $A \vee \neg A$ (Regel vom ausgeschlossenen Dritten; *tertium non datur*)
3. $\neg(A \wedge \neg A)$ (Regel vom ausgeschlossenen Widerspruch)

Beweis. 1. Wie betrachten die Wahrheitstabellen dieser Aussage für die beiden möglichen Werte von A . Da die Spalten A und $\neg(\neg A)$ gleich sind, sind auch die Aussagen A und $\neg(\neg A)$ äquivalent:

A	$\neg A$	$\neg(\neg A)$
w	f	w
f	w	f

Für 2. und 3. stellt man fest, dass unabhängig vom Wahrheitswert von A die im Satz behauptete Aussage stets den Wert w hat:

A	$\neg A$	$A \vee \neg A$
w	f	w
f	w	w

A	$\neg A$	$A \wedge \neg A$	$\neg(A \wedge \neg A)$
w	f	f	w
f	w	f	w

□

Aber was ist überhaupt ein Beweis? Wenn Sie Mathematiker fragen, werden die meisten wohl zunächst so etwas murmeln wie „Ein Beweis ist die Herleitung von Aussagen aus als wahr angenommenen Aussagen mittels der Regeln der Aussagenlogik“. Ein Beweis ist somit eine mechanischen Anwendung von logischen Operationen aus unserem Baukasten. Die Praxis sieht anders aus: wollte man jeden Beweis auf diese Weise aufschreiben, würden Beweise nicht nur sehr lang, sondern auch so komplex, dass sie kaum noch nachzuvollziehen sind. Daher werden viele Abkürzungen akzeptiert, wie etwa das Auslassen von Argumenten, die man schon etliche Male ausgeführt hat. Damit gelangt man zu der weitverbreiteten Auffassung, ein Beweis sei, was die Mehrheit aller Mathematiker als Beweis akzeptiert (was schon ein wenig selbstreferentiell erscheint ...). Dahinter liegt jedoch wieder die erste Antwort verborgen: als Beweis wird nur akzeptiert, wovon man sicher ist, dass es prinzipiell auf eine Kette mechanischer logischer Operationen zurückgeführt werden kann. Instruktiv ist dabei die Geschichte des Kepler-Problems und des daraus entstandenen QED-Projekts.

Eine Aussage, die aus der Verknüpfung mehrerer Aussagen hervorgeht und für alle möglichen Wahrheitswerte der verknüpften Aussagen wahr ist, nennt man eine Tautologie. Die erste Aussage des obigen Satzes ist eine solche Tautologie; weitere Beispiele sind:

Satz 1.2. *Seien A, B, C Aussagen. Dann gilt:*

- (a) $A \vee B \Leftrightarrow B \vee A$ (Kommutativgesetz)
- (b) $A \wedge B \Leftrightarrow B \wedge A$ (Kommutativgesetz)
- (b) $(A \vee B) \vee C \Leftrightarrow A \vee (B \vee C)$ (Assoziativgesetz)
- (b) $(A \wedge B) \wedge C \Leftrightarrow A \wedge (B \wedge C)$ (Assoziativgesetz)
- (c) $A \vee (B \wedge C) \Leftrightarrow (A \vee B) \wedge (A \vee C)$ (Distributivgesetz)
- (c) $A \wedge (B \vee C) \Leftrightarrow (A \wedge B) \vee (A \wedge C)$ (Distributivgesetz)
- (d) $\neg(A \vee B) \Leftrightarrow \neg A \wedge \neg B$ (Regeln von de Morgan)
- (d) $\neg(A \wedge B) \Leftrightarrow \neg A \vee \neg B$ (Regeln von de Morgan)
- (e) $(A \Rightarrow B) \Leftrightarrow \neg A \vee B$
- (f) $\neg(A \Rightarrow B) \Leftrightarrow A \wedge \neg B$

(g) $(A \Rightarrow B) \Leftrightarrow \neg B \Rightarrow \neg A$ (Kontrapositionsgesetz)

Beweis. Alle Behauptungen sind leicht durch Wahrheitstafeln zu zeigen; wir führen das hier für (c) und (e) vor, der Rest bleibt als Übung. Aus der Gleichheit der letzten beiden Spalten der Tafel

A	B	C	$A \vee B$	$A \vee C$	$B \wedge C$	$A \wedge (B \vee C)$	$(A \vee B) \wedge (A \vee C)$
w	w	w	w	w	w	w	w
w	w	f	w	w	f	w	w
w	f	w	w	w	f	w	w
w	f	f	w	w	f	w	w
f	w	w	w	w	w	w	w
f	w	f	w	f	f	w	w
f	f	w	f	w	f	w	w
f	f	f	f	f	f	f	f

folgt das erste Distributivgesetz, das zweite sieht man analog. Die Tafel

A	B	$\neg A$	$A \Rightarrow B$	$\neg A \vee B$
w	w	f	w	w
w	f	f	f	f
f	w	w	w	w
f	f	w	w	w

zeigt (e). □

Aus Teil (g) des Satzes folgt, dass man die Behauptung „aus A folgt B “ beweisen kan, indem man die äquivalente Behauptung „ist B falsch, so auch A “ beweist.

Mit der Regel vom ausgeschlossenen Dritten ergibt sich auch die Zulässigkeit der Methode des **indirekten Beweises**, die man auch wie folgt formulieren kann:

$$(A \Rightarrow B) \Leftrightarrow ((A \wedge \neg B) \Rightarrow \neg A)$$

d.h. wenn wir die Aussage B aus der Aussage A folgern wollen, können wir dies tun, indem wir die Annahme „es gilt A , aber nicht B “ zum Widerspruch führen.

2. Prädikatenlogik

Oft will man Aussagen formulieren, in denen eine oder mehrere Variable vorkommen; einen solchen Satz nennt man eine **Aussageform**. Dabei versteht man unter einer Variablen eine Leerstelle in einem logischen Ausdruck, in die ein konkretes Objekt eingesetzt werden kann. Es dürfen aber nur solche Objekte eingesetzt werden, die die Aussageform zu einer sinnvollen Aussage machen; man sagt, die Objekte stammen aus einem zulässigen Objektbereich oder auch Grundbereich (dies werden wir mit Mengenlehre etwas präzisieren).

Zum Beispiel ist

$$A(x): x \text{ ist eine Primzahl}$$

eine Aussageform; $A(3)$ ist wahr und $A(4)$ ist falsch, hingegen ist $A(1,5)$ keine sinnvolle Aussage: der Grundbereich dieser Aussageform sind die natürlichen Zahlen.

Eine Aussageform $A(x)$ mit nur einer Variablen beschreibt eine Eigenschaft, die das einzusetzende Objekt x haben soll. Dadurch bekommt x ein sogenanntes **Prädikat**, weshalb man die Lehre von den Aussageformen auch **Prädikatenlogik** nennt.

Zu einer gegebenen Aussageform $A(x)$ kann man Aussagen bilden wie

$$\text{„für alle } x \text{ gilt } A(x)\text{“, geschrieben } \forall x A(x) \text{ oder } \forall_x A(x) \text{ oder } \bigwedge_x A(x)$$

$$\text{„es gibt ein } x, \text{ so dass } A(x) \text{ gilt“, geschrieben } \exists x A(x) \text{ oder } \exists_x A(x) \text{ oder } \bigvee_x A(x)$$

Da diese Zusätze „es existiert,“ und „für alle“ die Aussageform quantifizieren, nennt man $\forall$ den Allquantor und $\exists$ den Existenzquantor. Die Notation ist in der Literatur nicht einheitlich; häufiger findet man $\forall$ (ein auf den Kopf gestelltes „A“) und $\exists$ (ein gespiegeltes „E“); die Symbole $\bigwedge$ und

$\forall$ scheinen jedoch, in Anlehnung an $\wedge$ und $\forall$, passender. Wir werden der Mehrheit folgen und $\forall$ und $\exists$ verwenden. Oft sieht man auch das Symbol $\exists$ (den durchgestrichenen Existenzquantor) mit der Bedeutung „es existiert kein“.

Satz 2.1. Sei $A(x)$ eine Aussageform. Dann gilt

$$\begin{aligned}\neg(\forall x A(x)) &\Leftrightarrow \exists x \neg A(x) \\ \neg(\exists x A(x)) &\Leftrightarrow \forall x \neg A(x)\end{aligned}$$

(Also könnte man auf einen der beiden Quantoren verzichten ...)

Oft gebrauchte Schreibweisen sind

$$\begin{aligned}\forall x, y A(x, y) &\text{ statt } \forall x \forall y A(x, y) \\ (\forall x \text{ mit } B(x)) A(x) &\text{ statt } \forall x (B(x) \Rightarrow A(x)) \\ (\exists x \text{ mit } B(x)) A(x) &\text{ statt } \exists x (A(x) \wedge B(x))\end{aligned}$$

Ein typisches Beispiel für diese abkürzende Schreibweise ist

$$\forall \varepsilon > 0 \exists \delta > 0 (|x - y| < \delta \Rightarrow |f(x) - f(y)| < \varepsilon)$$

Auch bei quantifizierten Aussagen muss man eventuell Klammern setzen. Dabei folgt man der Konvention, dass Quantoren Vorrang vor Junktoren haben. Eine Variable, die unmittelbar nach einem Quantor auftaucht, ist durch diesen *gebunden*, andere Variablen sind *frei*. Die Bindung der Variablen endet nach der ersten Aussageform, es sei denn, man hat Klammern gesetzt. Zum Beispiel endet die Wirkung des Existenzquantors in

$$\exists x x < 0 \wedge \exists y \forall x xy = 0$$

mit der ersten Formel („ $x < 0$ “), in der zweiten Formel ist x wieder frei (für den zulässigen Objektbereich „Zahlen“ bedeutet der obige Ausdruck in Worten: „es gibt eine negative Zahl, und zu jeder Zahl gibt es eine weitere, so dass das Produkt der beiden Zahlen Null ist“). Nur Variablen, die frei sind, sind überall in der Formel als Namen für die jeweils gleichen Objekte zu verstehen.

Beispiel. Die Aussagen (i) $\forall x A(x) \Rightarrow B$ und (ii) $\forall x (A(x) \Rightarrow B)$ sind verschieden: die erste bedeutet „gilt für jedes x die Aussage $A(x)$, so gilt auch die Aussage B “, hingegen (ii) „für jedes x gilt: $A(x)$ impliziert B “.

Zur Übung überlege man sich:

$$\neg(\forall x (B(x) \Rightarrow A(x))) \Leftrightarrow \exists x (B(x) \wedge \neg A(x))$$

und

$$\neg(\exists x (B(x) \wedge A(x))) \Leftrightarrow \forall x (B(x) \Rightarrow \neg A(x))$$

Zum Schluss dieses Abschnitts noch einige Bemerkungen zum Beweisen. An Beweismethoden (von denen schon ein paar erwähnt wurden) stehen zur Verfügung:

1. Der **direkte Beweis**: Ist A eine Aussage, deren Wahrheit feststeht, und ist die Aussage „ $A \Rightarrow B$ “ ebenfalls wahr, so ist auch B wahr. (Diese Methode ist auch unter dem Namen *modus ponens* bekannt.)
2. Sind die Implikationen $A \Rightarrow B$ und $B \Rightarrow C$ wahr, so ist auch die Implikation $A \Rightarrow C$ wahr.
3. Das **Kontrapositionsgesetz**:

$$(A \Rightarrow B) \Leftrightarrow (\neg B \Rightarrow \neg A)$$

4. Der **indirekte Beweis**: Um die Implikation $A \Rightarrow B$ zu zeigen, genügt es, die Implikation $(A \wedge \neg B) \Rightarrow C$ für eine falsche Aussage C zu zeigen.

5. Der **Existenzbeweis**: Um eine Aussage der Form „ $\exists x A(x)$ “ zu beweisen, reicht es, ein einziges x anzugeben, so dass die Aussage $A(x)$ wahr ist. Das ist aber oft anspruchsvoller als es aussieht, denn man muss ein solches x erst einmal finden. Manchmal ist es viel einfacher, solche Existenzaussagen indirekt zu beweisen!

Beispiele. 1. „Es gibt eine rationale Zahl, die die Gleichung $x^2 + 5x - 14 = 0$ erfüllt“: Das ist offenbar eine wahre Aussage, und der Beweis besteht schlicht im Nachrechnen, dass zum Beispiel $x = 2$ eine solche Lösung ist: $2^2 + 5 \cdot 2 - 14 = 4 + 10 - 14 = 0$. Wie man diese Lösung gefunden hat, spielt keine Rolle!

2. Anders sieht es bei der Aussage „jede Lösung der Gleichung $x^2 + 5x - 14 = 0$ ist eine rationale Zahl“: Hier muss man alle Lösungen finden (es sind $x = 2$ und $x = -7$, und beide sind rationale Zahlen).

3. „ $\exists x \in \mathbb{Q} \quad x^2 = 2$ “: Wie soll man zeigen, dass es etwas nicht gibt? Mit einem Widerspruchsbeweis, also indem man die Annahme, es gäbe ein solches x , ad absurdum führt! Man nehme also an, es gäbe eine rationale Zahl $x = \frac{p}{q}$, deren Quadrat 2 ist. Dabei kann man davon ausgehen, dass p und q keinen gemeinsamen Teiler haben, denn ansonsten können wir den Bruch solange kürzen, bis dies der Fall ist. Es folgt

$$2 = \left(\frac{p}{q}\right)^2 = \frac{p^2}{q^2}, \quad \text{also} \quad p^2 = 2q^2.$$

Folglich ist p^2 gerade, und dann muss auch p gerade sein, etwa $p = 2r$. Aber dann ist $2q^2 = (2r)^2 = 4r^2$, also ist $q^2 = 2r^2$ gerade und damit auch q gerade. Also haben p und q beide den Teiler 2, was ausgeschlossen war, und man hat den gewünschten Widerspruch.

Es gibt einige typische Fehler, die beim Argumentieren immer wieder gemacht werden. Einer davon ist besonders schwer auszurotten, da er in den Schule geradezu vorexerziert wird, nämlich die „Umkehrung der Beweislast“: oft wird, um eine behauptete Gleichung (vermeintlich) zu beweisen, mit eben dieser Gleichung gestartet, um diese dann so lange umzuformen, bis man zu einer offensichtlich richtigen Gleichung wie etwa „ $0 = 0$ “ gelangt. Umgekehrt muss man es machen! (Wie man einen Beweis aufschreibt hat nur marginal etwas damit zu tun, wie man darauf gekommen ist.)

Beispiel. Behauptung: für alle reellen Zahlen x und a gilt die Gleichung

$$\sqrt{x^2 - 2ax + a^2} + a = x.$$

„Beweis“: Ziehe auf beiden Seiten der Gleichung a ab:

$$\sqrt{x^2 - 2ax + a^2} = x - a$$

und quadriere beide Seiten:

$$x^2 - 2ax + a^2 = (x - a)^2$$

Die rechte Seite ist nach der zweiten binomischen Formel gleich der rechten, also stimmt die Behauptung. Aber was ist für $x = -a$? Da steht links $\sqrt{a^2 - 2a \cdot a + a^2} + a = \sqrt{0} + a = a$, aber rechts steht $-a$. Ist also stets $a = -a$? Das ist offensichtlich Unsinn; in dem „Beweis“ ist also etwas faul: Quadrieren ist keine Äquivalenzumformung, d.h. aus $p^2 = q^2$ folgt nicht $p = q$. Daher ist die Ausgangsgleichung nicht äquivalent zur offenbar für alle x und a wahren letzten Gleichung, sondern impliziert sie nur.

Ebenso gravierend ist der leider ebenso häufige „Beweis durch Beispiel“. Etwas subtiler, aber genauso falsch ist der „Beweis durch fehlendes Gegenbeispiel“, der immerhin so notorisch ist, dass er es in jede gängige Liste von Fehlschlüssen geschafft hat, unter dem schönen Namen *argumentum ad ignorantiam*. Auch nicht zu empfehlen ist die *petitio principii*, die Vorwegnahme des Grundes, d.h. die Inanspruchnahme eines Arguments, das selber des Beweises bedarf (mit dem Zirkelschluss als besonders häufiger Variante), oder die *ignoratio elenchi*, die irrelevante Schlussfolgerung.

2 Elementare Mengenlehre

1. Mengen

Cantor¹ versuchte 1895, den Mengenbegriff so zu fassen:

Definition 1.1. *Eine Menge ist eine Zusammenfassung von bestimmten, wohlunterschiedenen Objekten unserer Anschauung oder unseres Denkens zu einem Ganzen.*

Diese Mengendefinition soll hier zugrunde gelegt werden, obwohl sie nicht ohne Probleme ist, denn sie lässt Widersprüche zu, wie die Russellsche Antinomie: Ein Barbier ist jemand, der alle rasiert, die sich nicht selbst rasieren. (Aber wer rasiert den Barbier?) Aus diesem Grund wurde die axiomatische Mengenlehre entwickelt, d.h. man postuliert die Existenz gewisser Objekte, die man Mengen nennt, und stellt dann Regeln auf, mittels derer man Mengen aus Mengen konstruieren kann.

Die Objekte, die zu einer Menge M zusammengefaßt werden, nennt man die *Elemente* von M . Wir schreiben

$$\begin{aligned}x \in M & \quad \text{für „}x \text{ ist ein Element von } M\text{“,} \\x \notin M & \quad \text{für } \neg(x \in M), \text{ d.h. „}x \text{ ist kein Element von } M\text{“}\end{aligned}$$

Beispiel. Die folgenden Zahlenmengen werden sehr oft gebraucht und haben daher Standardbezeichnungen:

$\mathbb{N}$ ist die Menge der natürlichen Zahlen (einschließlich der Null),
 $\mathbb{Z}$ ist die Menge der ganzen Zahlen,
 $\mathbb{Q}$ ist die Menge der rationalen Zahlen,
 $\mathbb{R}$ ist die Menge der reellen Zahlen,

(Auf die Zahlenmengen soll später noch einmal eingegangen werden, aber im Augenblick setzten wir eine gewisse Kenntnis von ihnen voraus.)

Manche Mengen kann man durch Auflistung ihrer Elemente beschreiben:

$$A = \{a, b, c\}, \quad B = \{1, 17, 33, 987\}, \quad C = \{\text{rot, blau, grün, gelb, schwarz}\},$$

jedenfalls wenn sie nur endlich viele Elemente hat, oder auch mittels der berühmt-berüchtigten „Pünktchen“-Notation, wie etwa

$$M = \{1, 2, 3, \dots, 99\} \quad N = \{\dots, -3, -1, 1, 3, 5, \dots\}, \quad P = \{2, 3, 5, 7, 11, 13, \dots\}.$$

Dabei geht man davon aus, dass die Leser aus den (endlich vielen) angegebenen Elementen schon sehen können, welche weitere Elemente noch dazugehören. Im Fall der Mengen M und N ist eine Fehlinterpretation eher unwahrscheinlich, aber dass es sich bei der Menge P um die Menge aller Primzahlen handeln soll, ist weniger klar. Diese Unsicherheit kann man meist umgehen, wenn man die in Frage stehende Menge durch eine Eigenschaft, die ihre Elemente erfüllen sollen, beschreibt, also in der Form

$$D = \{x \in G \mid A(x)\}$$

¹Georg Cantor, deutscher Mathematiker, 1845–1918. Gemeinhin als der Begründer der Mengenlehre angesehen.

für eine Aussageform $A(x)$ und eine Grundmenge G . So ist in den obigen Beispielen

$$M = \{n \in \mathbb{N} \mid 0 < n < 100\}$$

$$N = \{n \in \mathbb{Z} \mid \exists m \in \mathbb{Z} \quad 2m - 1 = n\}$$

$$P = \{p \in \mathbb{N} \mid \forall m, n \in \mathbb{N} \ ((mn = p) \Rightarrow (n = 1 \vee m = 1))\}$$

Es gibt genau eine Menge, die keine Elemente enthält, die leere Menge $\emptyset$ (oder auch $\{\}$).

Definition 1.2. Eine Menge N heißt eine Teilmenge der Menge M , wenn jedes Element von N auch Element von M ist. Man schreibt dann $N \subset M$ oder $M \supset N$.

Zwei Mengen M, N sind gleich, geschrieben $M = N$, wenn sie die gleichen Elemente haben, also

$$\begin{aligned} M = N &\Leftrightarrow (M \subset N \wedge N \subset M) \\ &\Leftrightarrow ((x \in M \Rightarrow x \in N) \wedge (x \in N \Rightarrow x \in M)). \end{aligned}$$

(Dies ist der sogenannte extensionalistische Standpunkt der Mengenlehre.)

Aus gegebenen Mengen kann man durch bestimmte Operationen neue Mengen bilden:

1. durch Aussonderung, das heißt durch Angabe einer Bedingung: ist M eine Menge und $A(x)$ eine Aussageform, so ist

$$\{x \in M \mid A(x)\}$$

wieder eine Menge (siehe oben).

2. Die Vereinigung zweier Mengen A, B ist die Menge

$$A \cup B = \{x \mid x \in A \text{ oder } x \in B\}.$$

3. Der Durchschnitt oder Schnitt zweier Mengen A, B ist die Menge

$$A \cap B = \{x \mid x \in A \text{ und } x \in B\}.$$

4. Die Differenz zweier Mengen A, B ist die Menge

$$A - B = \{x \in A \mid x \notin B\},$$

manchmal auch geschrieben als $A \setminus B$. Ist B eine Teilmenge von A , so nennt man $A - B$ auch das *Komplement* von B in A .

5. Die symmetrische Differenz zweier Mengen A, B ist die Menge

$$A \Delta B = (A \cup B) - (A \cap B).$$

6. Die Potenzmenge einer Menge M ist die Menge aller Teilmengen von M :

$$\mathcal{P}(M) = \{T \mid T \subset M\}.$$

7. Das kartesische Produkt der Mengen A, B ist die Menge

$$A \times B = \{(x, y) \mid x \in A \text{ und } y \in B\}$$

sie ist die Menge aller geordneten Paare von Elementen von A bzw. B . Ein geordnetes Paar läßt sich auch schreiben als $(x, y) = \{\{x\}, \{x, y\}\}$.

Mengen können also wieder Mengen als Elemente haben, wie zum Beispiel die Potenzmenge einer Menge, oder sind M und N Mengen, so ist $\{M, N\}$ wieder eine Menge (diesen Prozess nennt man *Paarbildung*). Zu beachten ist dabei: $M \neq \{M\}$ (!).

Lemma 1.3 (Rechenregeln für Mengen). Seien A, B und C Mengen. Dann gilt

$$\begin{aligned} \text{(a)} \quad A \cup B &= B \cup A \\ A \cap B &= B \cap A \end{aligned} \quad (\text{Kommutativgesetz})$$

- (b) $(A \cup B) \cup C = A \cup (B \cup C)$ (Assoziativgesetz)
 $(A \cap B) \cap C = A \cap (B \cap C)$
(c) $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$ (Distributivgesetz)
 $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$
(d) $A - (B \cup C) = (A - B) \cap (A - C)$ (De Morgan)
 $A - (B \cap C) = (A - B) \cup (A - C)$

Beweis. Alle diese Behauptungen folgen aus den entsprechenden Regeln der Aussagenlogik. (a) und (b) sind klar. Zu (c): Es gilt

$$\begin{aligned}
 x \in A \cup (B \cap C) &\Leftrightarrow x \in A \vee (x \in B \wedge x \in C) \\
 &\Leftrightarrow (x \in A \vee x \in B) \wedge (x \in A \vee x \in C) \\
 &\Leftrightarrow x \in (A \cup B) \wedge x \in (A \cup C) \\
 &\Leftrightarrow x \in (A \cup B) \cap (A \cup C)
 \end{aligned}$$

Damit ist gezeigt, die beiden Mengen $A \cup (B \cap C)$ und $(A \cup B) \cap (A \cup C)$ dieselben Elemente haben, also gleich sind. Die zweite Aussage verifiziert man durch analoge Argumentation. Die erste der beiden De Morganschen Regeln sieht man wie folgt:

$$\begin{aligned}
 x \in A - (B \cup C) &\Leftrightarrow x \in A \wedge x \notin (B \cup C) \\
 &\Leftrightarrow x \in A \wedge \neg(x \in (B \cup C)) \\
 &\Leftrightarrow x \in A \wedge \neg(x \in B \vee x \in C) \\
 &\Leftrightarrow x \in A \wedge (\neg(x \in B) \wedge \neg(x \in C)) \\
 &\Leftrightarrow (x \in A \wedge \neg(x \in B)) \wedge (x \in A \wedge \neg(x \in C)) \\
 &\Leftrightarrow (x \in A \wedge x \notin B) \wedge (x \in A \wedge x \notin C) \\
 &\Leftrightarrow (x \in A - B) \wedge (x \in A - C) \\
 &\Leftrightarrow x \in (A - B) \cap (A - C)
 \end{aligned}$$

Die zweite Identität sieht man wieder analog. □

2. Abbildungen

Definition 2.1. Eine Abbildung oder Funktion f zwischen zwei Mengen M und N ist eine Vorschrift, die jedem Element $x \in M$ ein eindeutig bestimmtes Element $f(x) \in N$ zuordnet.

Man schreibt dann $f: M \rightarrow N$ für diese Abbildung und nennt M den Definitionsbereich sowie N den Wertebereich von f . Häufig benutzt man auch die Notation $x \mapsto f(x)$, um die Zuordnung auf dem Niveau der Elemente zu beschreiben. Zur Angabe einer Abbildung gehört außer der Zuordnungsvorschrift auch die Angabe des Definitionsbereichs M und des Wertebereichs N .

Oft lassen sich Funktionen in geschlossener Form angeben, etwa durch eine Rechenvorschrift, aber beileibe nicht immer.

Beispiele. 1. Durch die Zuordnungsvorschrift $x \mapsto x^2$ sind Abbildungen

$$f: \mathbb{R} \rightarrow \mathbb{R}, \quad g: \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}, \quad h: \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$$

definiert; diese Abbildungen sind verschieden.

2. Eine besonders einfache, aber wichtige Abbildung ist die identische Abbildung

$$\text{id}_M: M \rightarrow M \quad \text{mit} \quad x \mapsto x.$$

3. Sei $L(t)$ der Lärmpegel einer Baustelle zum Zeitpunkt t . Die Zuordnung $t \mapsto L(t)$ ist eine Funktion, die sicherlich nicht ohne weiteres durch eine Formel beschrieben werden kann.

4. Seien M und N Mengen. Die Abbildungen

$$\begin{aligned} \text{pr}_1: M \times N &\longrightarrow M & \text{mit } (x, y) &\mapsto x \\ \text{pr}_2: M \times N &\longrightarrow N & \text{mit } (x, y) &\mapsto y \end{aligned}$$

heißen die **kanonischen Projektionen**.

Sei nun $f: M \rightarrow N$ eine Funktion. Für eine Teilmenge $A \subset M$ nennt man die Menge

$$\text{Bild}(A) := \{f(x) \in N \mid x \in A\}$$

das **Bild** von A . Statt $\text{Bild}(M)$ schreibt man auch $\text{Bild}(f)$ und nennt diese Menge das Bild von f . Das **Urbild** einer Teilmenge $B \subset N$ ist die Menge

$$f^{-1}(B) := \{x \in M \mid f(x) \in B\};$$

diese Menge kann auch leer sein. Schließlich nennt man

$$\Gamma_f := \{(x, f(x)) \mid x \in M\} \subset M \times N$$

den **Graphen** der Funktion f .

Man prüft leicht nach:

Lemma 2.2. *Sei $f: M \rightarrow N$ eine Abbildung. Seien $A, B \subset M$ und $C, D \subset N$ Teilmengen. Dann gilt:*

- (a) $f(A \cup B) = f(A) \cup f(B)$
- (b) $f(A \cap B) \subset f(A) \cap f(B)$
- (c) $f^{-1}(C \cup D) = f^{-1}(C) \cup f^{-1}(D)$
- (d) $f^{-1}(C \cap D) = f^{-1}(C) \cap f^{-1}(D)$
- (e) $f(f^{-1}(C)) \subset C$.
- (f) $A \subset f^{-1}(f(A))$.

Übungsaufgabe. Man überlege sich durch Beispiele, dass in (b), (e) und (f) Gleichheit im allgemeinen nicht gelten kann.

Definition 2.3. *Seien $f: A \rightarrow B$ und $g: B \rightarrow C$ Abbildungen, so ist ihre **Komposition** (oder Verkettung) die Abbildung*

$$g \circ f: A \longrightarrow C \quad \text{mit} \quad (g \circ f)(x) = g(f(x)).$$

Für eine Abbildung $f: M \rightarrow N$ gilt offenbar

$$\text{id}_N \circ f = f = f \circ \text{id}_M.$$

Weiterhin ist die Komposition von Abbildungen assoziativ, d.h. sind $f: A \rightarrow B$, $g: B \rightarrow C$ und $h: C \rightarrow D$ Abbildungen, so dass die Kompositionen $g \circ f$ und $h \circ g$ erklärt sind, so gilt

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

Hingegen ist die umgekehrte Komposition $f \circ g$ im allgemeinen nicht definiert (es sei denn, $A = C$).

Definition 2.4. *Eine Abbildung $f: M \rightarrow N$ heißt*

- (a) *injektiv (oder eineindeutig), wenn aus $f(x) = f(y)$ schon $x = y$ folgt,*
- (b) *surjektiv, wenn es zu jedem $y \in N$ ein $x \in M$ gibt mit $f(x) = y$,*
- (c) *bijektiv, wenn sie injektiv und surjektiv ist.*

Eine bijektive Abbildung nennt man auch eine Bijektion, usw.

Eine Abbildung $f: M \rightarrow N$ ist also genau dann injektiv, wenn die Urbildmenge $f^{-1}(y)$ jedes Elements $y \in N$ höchstens ein Element hat, und surjektiv genau dann, wenn $f(M) = N$ ist.

Beispiele. 1. Die Abbildung $f: \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$ ist weder injektiv noch surjektiv, denn es gilt $f(-1) = f(1)$, und $-1 \in \mathbb{R}$ hat kein Urbild. Hingegen ist die Abbildung $g: \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}$ mit $g(x) = x^2$ surjektiv (denn jede nichtnegative Zahl hat eine Quadratwurzel), aber nach wie vor nicht injektiv. Die Abbildung $h: \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$ mit $h(x) = x^2$ schließlich ist bijektiv. Es kommt also nicht nur auf die Abbildungsvorschrift an!

2. Die Abbildung $f: \mathbb{R} \rightarrow \mathbb{R}$ gegeben durch

$$f(x) = \begin{cases} 0 & \text{falls } x = 0, \\ \frac{1}{x} & \text{falls } x \neq 0 \end{cases}$$

ist bijektiv.

Eine andere Charakterisierung injektiver und surjektiver Abbildungen ist:

Lemma 2.5. *Sei $f: M \rightarrow N$ eine Abbildung.*

- (a) *f ist genau dann injektiv, wenn es eine zu f **linksinverse** Abbildung gibt, d.h. eine Abbildung $g: N \rightarrow M$ gibt mit $g \circ f = \text{id}_M$.*
- (b) *f ist genau dann surjektiv, wenn es eine zu f **rechtsinverse** Abbildung gibt, d.h. eine Abbildung $g: N \rightarrow M$ mit $f \circ g = \text{id}_N$.*
- (c) *f ist genau dann bijektiv, wenn es eine zu f **inverse** Abbildung gibt, d.h. eine Abbildung $g: N \rightarrow M$ mit $g \circ f = \text{id}_M$ und $f \circ g = \text{id}_N$. Die inverse Abbildung ist eindeutig bestimmt.*

Man schreibt dann oft $f^{-1}: f(M) \rightarrow M$ für die auch Umkehrabbildung oder Umkehrfunktion genannte inverse Abbildung, aber man darf die Umkehrabbildung nicht mit der Urbildmenge verwechseln, obwohl dasselbe Symbol f^{-1} verwendet wird!

Beweis. (a) Sei f injektiv; zu konstruieren ist eine Linksinverse g . Sei ein Element $x_0 \in M$ fest gewählt und sei $y \in N$. Wenn es ein $x \in M$ gibt mit $f(x) = y$, so ist dieses x eindeutig bestimmt, da f injektiv ist. Setze also

$$g(y) := \begin{cases} x & \text{falls } f(x) = y, \\ x_0 & \text{falls } y \notin \text{Bild}(f). \end{cases}$$

Dann gilt offenbar $g(f(x)) = x$ für jedes $x \in M$. Umgekehrt: Gibt es eine rechtsinverse Abbildung g und sind $x, x' \in M$ mit $f(x) = f(x')$, so folgt $x = g(f(x)) = g(f(x')) = x'$.

(b) Sei f surjektiv; zu konstruieren ist eine Rechtsinverse g . Nach Voraussetzung gibt es zu jedem $y \in N$ ein Urbild x_y ; setze $g(y) = x_y$. Dann gilt $f(g(y)) = f(x_y) = y$. Gibt es andererseits eine Linksinverse g , so ist $g(y)$ ein Urbild von y unter f .

(c) Ist f bijektiv, so hat f nach (a) ein Linksinverses g_ℓ und nach (b) ein Rechtsinverses g_r . Folglich gilt

$$g_r \circ f = \text{id}_M \circ (g_r \circ f) = (g_\ell \circ f) \circ (g_r \circ f) = g_\ell \circ ((f \circ g_r) \circ f) = g_\ell \circ (\text{id}_N \circ f) = g_\ell \circ f = \text{id}_M,$$

d.h. g_r ist auch linksinvers zu f und damit invers zu f . Die andere Richtung folgt aus (a) und (b). Die Eindeutigkeit der Inversen schließlich sieht man so: Angenommen g und h seien beide invers zu f . Dann gilt

$$h = h \circ \text{id}_N = h \circ (f \circ g) = (h \circ f) \circ g = \text{id}_M \circ g = g.$$

□

Mittels Abbildungen lassen sich Mengen vergleichen; insbesondere kann man einen Begriff von ihrer Größe entwickeln:

Definition 2.6. *Seien M und N zwei Mengen. Wir nennen M und N gleichmächtig, wenn es eine bijektive Abbildung $\varphi: M \rightarrow N$ gibt. Eine Menge M heißt endlich, wenn es ein $n \in \mathbb{N}$ gibt und eine Bijektion $f: \{1, \dots, n\} \rightarrow M$, ansonsten heißt M unendlich.*

Ist M eine endliche Menge und $\varphi: \{1, \dots, n\} \rightarrow M$ eine Bijektion, so nennen wir n die Mächtigkeit von M und schreiben $n = \#M$. Die leere Menge $\emptyset$ hat keine Elemente, daher setzt man $\#\emptyset = 0$.

Satz 2.7. Seien A, B endliche Mengen. Dann gilt $\#(A \cup B) = \#A + \#B - \#(A \cap B)$.

Beweis. Die Aussage ist klar, falls eine der beiden Mengen leer ist. Sei also $A = \{a_1, \dots, a_n\}$ und $B = \{b_1, \dots, b_m\}$.

1. Sei $A \cap B = \emptyset$. Dann definiert

$$\varphi: \{1, \dots, m+n\} \longrightarrow A \cup B, \quad \varphi(i) = \begin{cases} a_i & \text{für } 1 \leq i \leq n, \\ b_{i-n} & \text{für } n+1 \leq i \leq n+m \end{cases}$$

eine bijektive Abbildung.

2. Sei $C = B - (A \cap B)$. Dann gilt $C \cap (A \cap B) = \emptyset$ und $B = C \cup (A \cap B)$, also nach 1. $\#B = \#C + \#(A \cap B)$. Weiterhin ist $C \cap A = \emptyset$ und $A \cup B = A \cup C$ und wiederum nach 1. gilt $\#(A \cup B) = \#A + \#C$. Insgesamt folgt

$$\#(A \cup B) = \#A + \#C = \#A + \#B - \#(A \cap B).$$

□

Satz 2.8. Seien X und Y endliche Mengen mit $\#X = \#Y = n$ und $f: X \rightarrow Y$ eine Abbildung. Dann sind äquivalent:

- (i) f ist injektiv,
- (ii) f ist surjektiv,
- (iii) f ist bijektiv.

Beweis. Es reicht, die Äquivalenz von (i) und (ii) nachzuweisen, der Rest ist Definition. Sei $X = \{x_1, \dots, x_n\}$.

(i) $\implies$ (ii): Da f injektiv ist, muß das Bild von f genau n Elemente haben, also schon ganz Y sein.

(ii) $\implies$ (i): Da f surjektiv ist, gilt $Y = \{f(x_1), \dots, f(x_n)\}$ und wegen $\#Y = n$ folgt, daß die Elemente $f(x_i)$ paarweise verschieden sein müssen. Aber dann ist f injektiv. □

Die Aussage des Satzes wird falsch, wenn wir die Endlichkeitsvoraussetzung fallen lassen. So sind zum Beispiel $\mathbb{N}$ und $\mathbb{Z}$ gleichmächtig, denn die Zuordnung

$$n \longmapsto \begin{cases} \frac{n}{2} & \text{falls } n \text{ gerade ist,} \\ \frac{-1-n}{2} & \text{falls } n \text{ ungerade ist} \end{cases}$$

definiert eine Bijektion $\mathbb{N} \rightarrow \mathbb{Z}$. Andererseits ist die Inklusion $\mathbb{N} \subset \mathbb{Z}$ eine injektive aber nicht surjektive Abbildung.

Ein weiteres schönes Resultat besagt, dass eine Menge nicht zu ihrer Potenzmenge gleichmächtig sein kann:

Satz 2.9 (Schröder-Bernstein). Ist M eine Menge, so kann keine surjektive Abbildung

$$f: M \rightarrow \mathcal{P}(M)$$

geben.

Beweis. Angenommen, es gäbe doch eine solche Abbildung $f: M \rightarrow \mathcal{P}(M)$. Sei dann

$$T := \{x \in M \mid x \notin f(x)\}.$$

Dies ist eine sinnvolle Definition. Da nun f surjektiv sein soll, muss es ein $m \in M$ geben mit $f(m) = T$. Aber das geht nicht: wäre $m \in T$, so müsste $m \notin f(m) = T$ sein, was unsinnig ist. Falls aber $m \notin T$ gilt, so ist $m \notin f(m)$, also ist m nach Definition von T doch ein Element von T , was wieder Unsinn ist. Es ist also weder $m \in T$ noch $m \notin T$ möglich. Dieser Widerspruch zeigt, dass die Menge T kein Urbild unter f haben kann. □

Das hier verwendete Beweisprinzip nennt man das *Cantorsche Diagonalverfahren*.

Es ist oft nützlich, Elemente einer Menge durch Indizes zu kennzeichnen. Dabei gehören die Indizes wiederum einer Menge an, der sogenannten *Indexmenge*. Formal geschieht dabei folgendes: Sei $f: I \rightarrow M$ eine Abbildung. Dann sagt man, die Menge $\text{Bild}(f)$ wird durch die Menge indiziert; anstelle von $f(i)$ schreibt man häufig f_i und spricht auch von der *Familie* $(f_i)_{i \in I}$ oder $\{f_i\}_{i \in I}$. Das Element f_i heißt dann die i -te Komponente der Familie. Solche Familien heißen auch I -Tupel (von Elementen von M). Ist $I = \{1, 2, \dots, n\}$, spricht man gemeinhin von n -Tupeln. 2-Tupel sind dann geordnete Paare (m_1, m_2) , 3-Tupel sind Tripel (m_1, m_2, m_3) , usw.

Die Mengen-Operationen Vereinigung, Schnitt und Produkt lassen sich auch auf Familien ausdehnen: sei I eine nichtleere Indexmenge und $(A_i)_{i \in I}$ eine Familie von Mengen. Dann ist

$$\bigcup_{i \in I} A_i = \{x \mid x \in A_i \text{ für wenigstens ein } i\}$$

die Vereinigung und

$$\bigcap_{i \in I} A_i = \{x \mid x \in A_i \text{ für alle } i \in I\}$$

der Durchschnitt der Mengen A_i , $i \in I$. Für die endliche Indexmenge $I = \{1, 2, \dots, n\}$ schreibt man auch

$$A_1 \cup \dots \cup A_n = \bigcup_{i=1}^n A_i \quad \text{bzw.} \quad A_1 \cap \dots \cap A_n = \bigcap_{i=1}^n A_i$$

für Vereinigung bzw. Schnitt.

Die Menge der I -Tupel $(a_i)_{i \in I}$ von Elementen aus $A := \bigcup_{i \in I} A_i$ mit $a_i \in A_i$ ist das kartesische Produkt der Mengen A_i , bezeichnet mit

$$\prod_{i \in I} A_i.$$

Für $I = \{1, \dots, n\}$ ist auch

$$A_1 \times \dots \times A_n = \prod_{i=1}^n A_i$$

gebräuchlich. Für ein $j \in I$ heißt die Abbildung

$$\text{pr}_j: \prod_{i \in I} A_i \longrightarrow A_j, \quad (a_i)_{i \in I} \mapsto a_j$$

die j -te Projektion.

3. Relationen

Definition 3.1. Eine Relation zwischen den Mengen A und B ist eine Teilmenge $R \subset A \times B$. Bei $A = B$ spricht man auch von einer Relation auf der Menge A .

Ist R eine Relation zwischen A und B und ist $(a, b) \in R$, so schreibt man auch aRb .

Beispiele. 1. Sei A eine Menge. Gleichheit ist eine Relation, die durch die *Diagonale*

$$\Delta_A := \{(a, a) \mid a \in A\} \subset A \times A$$

gegeben ist.

2. Für eine Funktion $f: A \rightarrow B$ definiert ihr Graph $\Gamma_f \subset A \times B$ eine Relation; diese hat die Eigenschaft, dass es zu $x \in A$ genau ein $y \in B$ gibt mit xRy (nämlich $y = f(x)$).

Definition 3.2. Sei A eine Menge. Eine Relation R auf A heißt **Ordnungsrelation**, wenn gilt:

- (i) Für alle $a \in A$ gilt aRa (**Reflexivität**)
- (ii) Für alle $a, b \in A$ gilt: $aRb \wedge bRa \Rightarrow a = b$ (**Antisymmetrie**)
- (iii) Für alle $a, b, c \in A$ gilt: $aRb \wedge bRc \Rightarrow aRc$ (**Transitivität**)

Gilt zusätzlich noch

- (iv) Für alle $a, b \in A$ gilt aRb oder bRa ,

so heißt R eine **Totalordnung** auf A .

Ordnungsrelationen werden meist mit dem Symbol $\leq$ bedacht. Man schreibt dann statt „ $a \leq b$ und $a \neq b$ “ kürzer $a < b$ sowie $a \geq b$ anstelle von $b \leq a$. Eine geordnete Menge ist ein Paar $(A, \leq)$ bestehend aus einer Menge A und einer Ordnungsrelation $\leq$ auf A ; die Menge heißt total geordnet, wenn $\leq$ eine Totalordnung ist. Man nennt zwei Elemente a, b einer geordneten Menge A **vergleichbar**, wenn $a \leq b$ oder $b \leq a$ gilt. (Aber nicht alle Elemente müssen vergleichbar sein.)

Beispiele. 1. Die natürliche Ordnung $\leq$ ist eine Totalordnung auf $\mathbb{R}$ und damit auch auf jeder Teilmenge von $\mathbb{R}$.

2. Sei M eine Menge. Dann definiert die Inklusion eine Ordnung auf der Potenzmenge $\mathcal{P}(M)$, die keine Totalordnung ist.

3. Sei $(A_i, \leq)_{i \in I}$ eine Familie geordneter Mengen. Dann definiert die Vorschrift

$$(a_i) \leq (b_i) \Leftrightarrow a_i \leq b_i \text{ für alle } i$$

eine Ordnung auf dem Produkt $\prod_i A_i$, die sogenannte Produktordnung. Selbst wenn jede der Mengen A_i total geordnet ist, ist die Produktordnung im allgemeinen keine Totalordnung.

4. Seien $A_1, \dots, A_n$ total geordnete Mengen. Definiere

$$(a_1, \dots, a_n) \leq_{\text{lex}} (b_1, \dots, b_n) \Leftrightarrow \begin{array}{l} (a_1, \dots, a_n) = (b_1, \dots, b_n) \\ \text{oder} \\ a_i \leq b_i \text{ für das kleinste } i \text{ mit } a_i \neq b_i. \end{array}$$

Diese Ordnung ist eine Totalordnung und heißt (aus ziemlich offensichtlichen Gründen) die **lexikographische** Ordnung.

Ein Element a^* einer geordneten Menge $(A, \leq)$ heißt *maximal*, wenn $a \leq a^*$ für jedes mit a^* vergleichbare Element a gilt, und ein *größtes Element*, wenn $a \leq a^*$ für alle $a \in A$; entsprechend sind minimale und kleinste Elemente definiert. Schließlich nennt man eine Totalordnung auf einer Menge A eine **Wohlordnung**, wenn jede nichtleere Teilmenge ein kleinstes Element besitzt.

Eine zweite, wichtige Klasse von Relationen erlaubt das Identifizieren von an sich verschiedenen Objekten, die man als äquivalent ansehen möchte:

Definition 3.3. Sei A eine Menge. Eine Relation R auf A heißt **Äquivalenzrelation**, wenn gilt:

- (i) Für alle $a \in A$ gilt aRa (**Reflexivität**)
- (ii) Für alle $a, b \in A$ gilt: $aRb \Rightarrow bRa$ (**Symmetrie**)
- (iii) Für alle $a, b, c \in A$ gilt: $aRb \wedge bRc \Rightarrow aRc$ (**Transitivität**)

Für Äquivalenzrelationen benutzt man meist das Symbol $\sim$. Zwei Elemente a, b heißen äquivalent, wenn $a \sim b$ (und damit auch $b \sim a$) gilt. Die Menge der zu einem $a \in A$ äquivalenten Elemente nennt man die **Äquivalenzklasse** von a , oft bezeichnet mit $[a]$, gelegentlich auch mit $\bar{a}$. Es ist also

$$[a] = \{x \in A \mid x \sim a\}.$$

Da Äquivalenzrelationen symmetrisch ist, sind Äquivalenzklassen nicht leer, denn es ist $a \in [a]$.

Satz 3.4. Sei $\sim$ eine Äquivalenzrelation auf der Menge A und seien $a, b \in A$. Dann sind äquivalent:

- (i) $[a] \cap [b] \neq \emptyset$
- (ii) $a \sim b$
- (iii) $[a] = [b]$

Beweis. (i) $\Rightarrow$ (ii): Sei $c \in [a] \cap [b]$. Dann gilt $a \sim c$ und $c \sim b$, wegen Transitivität also auch $a \sim b$. (ii) $\Rightarrow$ (iii): Sei $x \in [a]$. Dann gilt $x \sim a$ und $a \sim b$, also $x \sim b$ und damit $x \in [b]$. Dies zeigt $[a] \subset [b]$; die andere Inklusion folgt auch, da das Argument symmetrisch ist. (iii) $\Rightarrow$ (i) ist klar. Damit ist der Ringschluss komplett. $\square$

Es folgt, dass eine Äquivalenzrelation auf einer Menge A diese Menge in paarweise disjunkte Teilmengen zerlegt.

Anhang: Axiome der Mengenlehre

Das nachfolgende Axiomensystem wird mit **ZFC** bezeichnet; ohne das 10. Axiom („Choice“) nur mit **ZF**, nach Zermelo und Fraenkel.

1. Existenz der leeren Menge:

Es gibt eine Menge $\emptyset$, so dass für jede Menge x gilt: $x \notin \emptyset$.

2. Extensionalität:

Zwei Mengen x, y sind genau dann gleich, wenn sie die gleichen Elemente enthalten, d.h. wenn gilt: $\alpha \in x \Leftrightarrow \alpha \in y$.

3. Paarbildung:

Seien x, y Mengen, so existiert eine Menge M , die genau x und y als Elemente hat. Man schreibt dann $M = \{x, y\}$.

4. Vereinigungsmenge:

Sei X eine Menge, dann ist auch $\bigcup X = \bigcup_{x \in X} x$ eine Menge.

5. Potenzmenge:

Sei x eine Menge, dann ist auch $\mathcal{P}(x) = \{y \subset x\}$ eine Menge.

6. Unendlichkeitsaxiom:

Es gibt eine Menge, die $\emptyset$ enthält und mit jedem Element x auch $x \cup \{x\}$.

7. Fundierungsaxiom:

Jede nichtleere Menge x enthält ein Element y mit $x \cap y = \emptyset$.

8. Aussonderung:

Sei X eine Menge und P ein einstelliges Prädikat (eine Eigenschaft). Dann ist $\{x \in X \mid P(x)\}$ eine Menge.

9. Ansammlung:

Sei X eine Menge und P ein zweistelliges Prädikat. Wenn es zu jedem $x \in X$ ein y gibt, so dass $P(x, y)$, so existiert eine Menge Y , aus der passende y gewählt werden können.

10. Auswahlaxiom:

Sei X eine Menge, so dass jedes Element nichtleer ist. Dann existiert eine Funktion f , die jedem $x \in X$ ein Element aus x zuordnet, d.h. $\forall x \in X \quad f(x) \in x$.

Nützlich ist noch der Begriff einer **Klasse** als einer Gesamtheit, die von einer Eigenschaft aus dem Mengenuniversum ausgesondert werden. (Klassen lassen also de Russelsche Antinomie zu!)

Mit dem Klassenbegriff kann man anstelle der Ansammlung auch folgendes Axiom verwenden:

9'. Ersetzung:

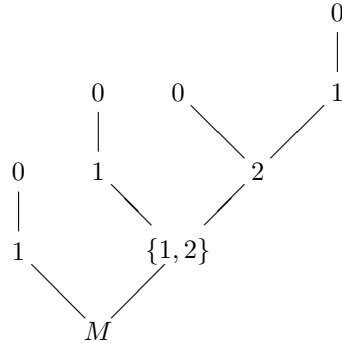
Ist F eine Funktion von einer Klasse D_F in die Klasse R_F , so ist das Bild jeder Menge X unter F wieder eine Menge.

Das Unendlichkeitsaxiom garantiert die Existenz einer Menge $\mathbb{N}$ mit den Elementen

$$\begin{aligned} 0 &= \emptyset, \\ 1 &= \{\emptyset\}, \\ 2 &= \{\emptyset, \{\emptyset\}\} = \{0, 1\} \\ 3 &= \{\emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\} = \{0, 1, 2\} \\ &\vdots \end{aligned}$$

Äquivalent ist zum Fundierungsaxiom folgende Formulierung. Betrachte den *Baum* der Menge M : die Wurzel ist M , darüber liegen alle Elemente von M , usw.. Dann hat jeder Zweig endliche Länge. (Das letzte Element nennt man dann ein *Blatt* des Baumes.)

Beispiel: $M = \{1, \{1, 2\}\}$ (wobei $0 = \emptyset$, $1 = \{\emptyset\}$, $2 = \{\emptyset, \{\emptyset\}\}$, usw.).



Denn: Angenommen es existiert ein unendlich langer Zweig. Dann gibt es eine Folge $m_0 = M$, $m_1 \in m_0, \dots, m_n \in m_{n-1}, \dots$. Aber dann ist $X = \{m_0, m_1, m_2, \dots\}$ eine nichtleere Menge, so dass jedes Element nichtleeren Schnitt mit X hat.

Bemerkungen. 1) Das Fundierungsaxiom macht die Russellsche Antinomie unmöglich: keine Menge ist Element ihrer selbst!

Beweis. Sei A eine Menge mit $A \in A$. Definiere $B = \{A\}$ (Paarbildungsaxiom). Aus dem Fundierungsaxiom folgt, dass A , das einzige Element von B , leeren Schnitt mit B haben muss, d.h. $A \cap B = \emptyset$. Aber $A \in A$ und $A \in B$ impliziert $A \in A \cap B$, ein Widerspruch. $\square$

2) Sei f eine Funktion auf $\mathbb{N}$ mit $f(n+1) \in f(n)$ für alle n . Definiere $S = \{f(n) : n \in \mathbb{N}\}$; dies ist eine Menge. Dann impliziert das Fundierungsaxiom (F): es existiert ein $f(k)$ mit $f(k) \cap S = \emptyset$. Aber per Definition von f und S ist $f(k+1) \in f(k) \cap S$, ein Widerspruch. Also folgt aus (F), dass es eine solche Funktion f nicht geben kann.

3) Bezeichne nun (F') das folgende (schwächere) Axiom:

(F') Es gibt keine unendlich absteigende Folge von Mengen.

Dann gilt: (F') + (Auswahlaxiom) $\Rightarrow$ (F).

Beweis. Sei S eine Menge, so dass für jedes $s \in S$ gilt: $s \cap S \neq \emptyset$. Sei g eine Auswahlfunktion für S , d.h. $g(s) \in s$. Definiere eine Funktion f auf $\mathbb{N}$ wie folgt:

$$\begin{aligned} f(0) &= g(S) \\ f(n+1) &= g(f(n) \cap S) \end{aligned}$$

Dann gilt $f(n) \in S$ für jedes n und $f(n+1) \in f(n)$. $\square$

4) Das Fundierungsaxiom erlaubt die Definition des geordneten Paares (a, b) als $\{a, \{a, b\}\}$ (statt $\{\{a\}, \{a, b\}\}$).

5) Das Auswahlaxiom ist äquivalent zum **Wohlordnungssatz**: *Jede Menge kann wohlgeordnet werden.*

Gödel² hat gezeigt: Es ist nicht möglich, die Widerspruchsfreiheit dieses Axiomensystems **ZFC** zu beweisen!

²Kurt Gödel, 1906–1978, einer der bedeutendsten Logiker des 20. Jahrhunderts. Seine berühmten „Unvollständigkeitssätze“ zeigten, dass Hilberts Programm, die Widerspruchsfreiheit der Mathematik zu nachzuweisen, zum Scheitern verurteilt war. Ein vollständiger Beweis des ersten Unvollständigkeitssatzes findet man auch in D. Hofstadters Buch „Gödel, Escher, Bach“.