

Einführung in Sprache und Grundbegriffe der Mathematik

Markus Junker
Mathematisches Institut
Albert–Ludwigs–Universität Freiburg

Wintersemester 2010/11

Inhaltsverzeichnis

0	Vorbemerkung	3
1	Mengen und elementare mathematische Strukturen	5
1.1	Naive Mengenlehre	5
1.2	Elementare Konstruktionen und ihre Eigenschaften	6
1.3	Relationen	10
1.4	Abbildungen	14
1.5	Unendliche Konstruktionen	17
2	Die natürlichen Zahlen und das Induktionsprinzip	21
2.1	Beweis per vollständiger Induktion	22
2.2	Summen und Produkte	23
3	Logik (oder: „Wie redet man in der Mathematik?“)	24
3.1	Logische Folgerung und logische Äquivalenz	25
3.2	Symbolik	27
3.3	Wichtige Regeln	29
3.4	Beweise und logische Schlüsse	29
3.5	Quantoren und Variable	31
3.6	Verneinungen	35
4	Mathematische Terminologie	37
4.1	Satz, Definition, Beweis	37
4.2	Weiterer mathematischer Sprachgebrauch	38
5	Anhang	40
5.1	Axiomatische Mengenlehre	40
5.2	Ein Exkurs über das Auswahlaxiom	40
5.3	Die Peano–Axiome	41

Vorbemerkung

Wozu dient dieses Skript? Die folgenden Seiten sollen Studierenden der Mathematik im ersten Semester eine Hilfe beim Erlernen der mathematischen Sprache bieten. Für den Idealfall, dass beide Grundvorlesungen (Analysis I und Lineare Algebra I) im ersten Semester gehört werden, kann es ergänzend als Nachschlagewerk genutzt werden. Für den Fall, dass nur eine der beiden Grundvorlesungen im ersten Semester gehört werden kann, können anhand dieses Kurzsriptes diejenigen Grundlagen vor- oder nachgearbeitet werden, die in der anderen Vorlesung eingeführt oder behandelt werden.



WARNUNG I



Dieses Skript stellt keinen Einführungskurs in die Mathematik dar und ist dafür auch nicht geeignet, sondern es ist dazu gedacht, bei Bedarf begleitend zu Vorlesungen und Übungen konsultiert werden. Es vor dem Studium durchzulesen ist ähnlich, wie wenn man vor dem Erlernen einer Fremdsprache Grammatiktafeln durcharbeitet: Eine trockene und abschreckende Angelegenheit, für die das wahre Verständnis fehlt. Man sollte insbesondere nicht glauben, dass man diese Seiten verstanden haben muss, bevor man mit dem Studium beginnen kann, sondern umgekehrt wird sich durch das Studium das Verständnis erschließen.

Was setzt es voraus? Das Skript setzt die Vertrautheit im Umgang mit mathematischer Symbolik und mathematischen Objekten voraus, die man in der Schule erworben haben sollte. Außerdem braucht man bei der Lektüre die Duldsamkeit, nicht alles beim ersten Lesen schon verstehen zu können, denn:



WARNUNG II



Dieses Skript benutzt selbst die mathematische Sprache und Denkweise, die es erklären will. Da die Anordnung systematisch ist, wird einiges erst später erklärt, was vorher schon verwendet wird. Somit sind eigentlich (mindestens) zwei Lekturedurchgänge erforderlich. Erfahrungsgemäß verschwinden eventuelle Schwierigkeiten aber weniger durch die Theorie als durch die mathematische Praxis (Vorlesung, Übungen).

Farbenlehre: In kleinerer, blauer Schrift sind weitergehende Anmerkungen angefügt, die übersprungen werden können. Diese Anmerkungen behandeln oft Sonderfälle, sind manchmal aber auch historischer oder philosophischer Natur. Begriffe, die eingeführt bzw. definiert werden, sind rot unterlegt; erläuternde Beispiele grün.

Auf den folgenden Seiten werden sehr viele mathematische Sachverhalte behauptet, insbesondere in den Listen mathematischer Gesetze, aber nicht bewiesen. Dies ist als Aufforderung gedacht, selbst über den Beweis nachzudenken. Allerdings wäre es sehr ermüdend, sämtliche Gesetze z. B. in Tabelle 1 auf Seite 9 nachzuprüfen. Stattdessen sollte man an einzelnen Beispielen die Art und Weise der Beweise einüben, und eventuell unplausibel wirkende Regeln nachrechnen.

Worüber redet eigentlich man in der Mathematik?

Es gibt eine Vielfalt an grundlegenden Objekten in der Mathematik (Zahlen, Punkte, Funktionen, Mengen ...). Sie sind keine sinnlich erfassbaren Dinge, sondern abstrakte Vorstellungen. Manchmal lassen sich diese Objekte (zumindest in einer Annäherung) visualisieren, zum Beispiel ein geometrisches Gebilde durch eine Zeichnung; man muss sich dann aber darüber im Klaren sein, dass die Zeichnung das Objekt nur abbildet und nicht mit ihm identisch ist.

Auch der Ursprung der mathematischen Objekte kann vielfältig sein: Teils ergeben sie sich als Abstraktionen aus dem Alltag, teils dienen sie der Modellierung, teils dem innermathematischen Verständnis. Mathematische Objekte werden durch ihre Funktionsweise verstanden, also durch ihre Eigenschaften und durch ihr Verhältnis zueinander. Sie werden daher oft auf axiomatische

Weise eingeführt: *Axiome* beschreiben, welche gewünschten Eigenschaften die Objekte haben sollen.

Ob man sinnvollerweise sagen kann, dass solche mathematischen Objekte existieren, und um welche Art von Existenz es sich dann handelt, ist ein philosophisches Problem, um das sich Mathematiker üblicherweise nicht kümmern. Ihnen reicht es, dass ihre Theorien „funktionieren“. Dies bedeutet, dass in den Theorien keine Widersprüche auftreten und dass sie geeignet sind, mathematische oder außer-mathematische Phänomene zu beschreiben.

Da die Widerspruchsfreiheit einer Theorie in der Regel aber nicht bewiesen werden kann, wird üblicherweise erwartet, dass man sie auf eine anerkannte Theorie zurückführt, zumeist auf das mengentheoretische Axiomensystem ZFC. Für mathematische Objekte bedeutet dies, dass man ihre „relative Existenz“ zeigt, d. h. dass man sie zum Beispiel mengentheoretisch modelliert: Man sucht also etwa in der durch ZFC beschriebenen Mengentheorie Objekte, die sich genau so verhalten, wie die Axiome für diese Objekte es vorschreiben. Aufgrund des jahrzehntelangen erfolgreichen Gebrauchs geht man allgemein davon aus, dass die durch ZFC beschriebene Mengenlehre widerspruchsfrei ist, oder zumindest, dass eventuell einmal auftretende Widersprüche nicht den Kern der Mathematik betreffen, sondern durch leichte Veränderungen des Axiomensystems ausgemerzt werden können.

Ebenso wichtig wie die Objekte der Mathematik sind die Bezüge zwischen ihnen. Diese können sich zum Beispiel durch Relationen ausdrücken (etwa die zweistellige Ordnungsrelation auf den natürlichen Zahlen – *eine Zahl ist größer als eine andere oder nicht* – oder die dreistellige „Zwischenrelation“ auf den Punkten einer Linie – *ein Punkt liegt zwischen zwei anderen oder nicht*) oder durch die Existenz von Operationen mit gewissen Eigenschaften (z. B. die Addition natürlicher Zahlen) oder durch kompliziertere Verhältnisse. Diese Relationen, Operationen etc. können dann selbst wieder als mathematische Objekte aufgefasst werden.

Eine Menge von Objekten mit gegebenen Bezügen zueinander wird üblicherweise eine *Struktur* genannt. Ein Großteil der mathematischen Arbeit besteht darin, Strukturen mit bestimmten Eigenschaften zu untersuchen und Sachverhalte darüber zu beweisen (*mathematische Sätze* oder *Theoreme* genannt). Insbesondere will man solche Strukturen klassifizieren und Strukturen finden, die als Modelle für Phänomene aus der Physik, der Wirtschaft etc. dienen können.

1 Mengen und elementare mathematische Strukturen

1.1 Naive Mengenlehre

Der Begriff der Menge ist einer der grundlegendsten Begriffe der Mathematik. Eine Menge kann man sich vorstellen als eine Zusammenfassung von verschiedenen Objekten (den *Elementen der Menge*) zu einem neuen Objekt. Diese Zusammenfassung soll so vor sich gehen, dass dabei das *Extensionalitätsprinzip* gilt:

Menge
Element

Zwei Mengen sind genau dann gleich, wenn sie die gleichen Elemente enthalten.

Eine Menge ist also eindeutig dadurch bestimmt, dass man weiß, welche Elemente sie enthält. Man schreibt $a \in M$ dafür, dass a ein Element von M ist (und sagt auch „ a liegt in M “), und $a \notin M$ dafür, dass a kein Element von M ist. Oft schreibt man abkürzend $a, b \in M$ und meint „ $a \in M$ und $b \in M$ “, und ähnlich für mehrere Elemente.

$\in, \notin$

Die *leere Menge* hat kein Element; nach dem Extensionalitätsprinzip ist sie durch diese Eigenschaft eindeutig bestimmt. Sie wird meist mit $\emptyset$ bezeichnet, manchmal auch mit $\{ \}$.

$\emptyset$

Man muss eine Menge von Objekten unterscheiden von einer Beschreibung dieser Menge. Wenn man eine Beschreibung hat (z. B. „die um die Sonne kreisenden Planeten“), so ist die *Extension* dieser Beschreibung gerade die Menge der die Beschreibung erfüllenden Objekte. Es kann durchaus sein, dass zwei Beschreibungen mit unterschiedlicher *Intension* die gleiche Extension haben, also (gewissermaßen zufällig) die gleiche Menge beschreiben, z. B. „die Monde der Venus“ und „die im Jahr 2010 in Deutschland herrschenden Könige“, die beide die leere Menge als Extension haben. Ein mathematisches Beispiel: „Die geraden Primzahlen“ und „die Zahlen, die halb so groß wie ihr Quadrat sind“ beschreiben beide die Menge $\{2\}$.

Eine endliche Menge kann gegeben sein durch eine Beschreibung ihrer Elemente oder durch die Auflistung ihrer Elemente innerhalb von geschweiften Klammern, wie zum Beispiel

$\{ \dots \}$

$\{10, 2, 3, 7, 5\}$

Wegen des Extensionalitätsprinzips spielt die Reihenfolge der Elemente keine Rolle und es können auch Elemente mehrfach genannt sein. Es gilt also z. B. $\{2, 1, 1, 3, 2\} = \{1, 2, 3\}$. Die Anzahl der Elemente einer endlichen Menge M wird meistens mit $|M|$ bezeichnet, es ist also $|\{5, 3, 5, 11\}| = 3$. Manche schreiben stattdessen $\#M$.

$|\dots|$

Unendliche Mengen kann man eigentlich nur durch Beschreibungen der Elemente angeben. Manchmal schreibt man auch den Anfang einer Auflistung, die mit drei Pünktchen endet. Dann muss aber klar sein, wofür die Pünktchen stehen. Als Beispiel folgen drei Möglichkeiten, die Menge der Primzahlen aufzuschreiben:

$$\{x \mid x \text{ ist natürliche Zahl und Primzahl}\} = \{x \in \mathbb{N} \mid x \text{ prim}\} = \{2, 3, 5, 7, 11, 13, 17, \dots\}$$

Es gibt einige Varianten dieser Schreibweisen; oft sieht man einen Doppelpunkt anstelle des senkrechten Striches.

Häufige Verwendung findet die „Pünktchenschreibweise“ auch in der Form $\{a_1, \dots, a_n\}$. Dies bedeutet: n ist eine Variable für eine beliebige natürliche Zahl und die Menge besteht aus mit a_i bezeichneten Elementen, wobei der Index i alle natürlichen Zahlen von 1 bis n durchläuft. Dabei ist es durchaus möglich, dass einige der a_i untereinander gleich sind. In Formeln ausgedrückt:

$a_1, \dots, a_n$

$$\{a_1, \dots, a_n\} = \{x \mid \text{es gibt ein } i \in \mathbb{N} \text{ mit } 1 \leq i \leq n, \text{ so dass } x = a_i\} \quad (*)$$

Insbesondere ist:

$$\begin{aligned} \{a_1, \dots, a_2\} &= \{a_1, a_2\} \\ \{a_1, \dots, a_3\} &= \{a_1, a_2, a_3\} \\ \{a_1, \dots, a_4\} &= \{a_1, a_2, a_3, a_4\} \quad \text{usw.} \end{aligned}$$

An vielen Stellen treten analoge Schreibweisen auf, ohne dass die Bedeutung jedesmal erneut erklärt wird.

Solche Schreibweisen sind für eine beliebige natürliche Zahl n gedacht, die sowohl sehr große als auch sehr kleine Werte annehmen kann. Die Schreibweise soll daher auch in „Extremsituationen“ funktionieren. Gewissermaßen ist der Fall $n = 2$ schon solch eine Extremsituation, weil die Pünktchen dann für nichts mehr stehen. Noch extremer sind die Fälle $n = 1$ und gegebenenfalls $n = 0$. Hierfür gilt:

$$\begin{aligned}\{a_1, \dots, a_1\} &= \{a_1\} \\ \{a_1, \dots, a_0\} &= \emptyset\end{aligned}$$

Man kann dies in der Erklärung (*) daraus ablesen, dass es im Falle $n = 1$ nur ein einziges solches i gibt, nämlich $i = 1$, und im Fall $n = 0$ gar kein solches i gibt. Man kann dies fürs erste aber auch als eine Konvention ansehen. (Dann muss man sich aber immer wieder davon überzeugen, dass es deshalb eine sinnvolle Konvention ist, weil alle Regeln, die für Mengen von der Form $\{a_1, \dots, a_n\}$ angegeben werden, auch für die Fälle $n = 1$ und $n = 0$ gelten.)

Eine Menge M heißt **Teilmenge** einer Menge N , wenn jedes Element von M auch Element von N ist, und sie heißt **echte Teilmenge**, wenn darüberhinaus $M \neq N$. Es gibt dafür zwei gleichermaßen verbreitete Symbolschreibweisen (von denen ich Version 1 benutzen werde):

Teilmenge
 $\subset, \subseteq, \subsetneq$

	Teilmenge	echte Teilmenge
Version 1:	$M \subseteq N$	$M \subset N$
Version 2:	$M \subset N$	$M \subsetneq N$ oder $M \subsetneqq N$

Mengen sind mathematische Objekte und können daher auch wieder zu Mengen zusammengefasst werden, d. h. es gibt Mengen, deren Elemente selbst wieder Mengen sind. Die einfachsten solchen Mengen sind von der folgenden Gestalt: Zu jeder Menge M gibt es die Menge $\{M\}$, die als einziges Element die Menge M enthält. Daraus kann man wiederum die Menge $\{\{M\}\}$ bilden, usw.

Dies geht auch für die leere Menge: Es gibt die Menge $\{\emptyset\}$, und dies ist eine andere Menge als die leere Menge selbst. Die leere Menge hat kein Element; die Menge $\{\emptyset\}$ hat ein Element, nämlich die leere Menge.

Am Anfang ist dies oft verwirrend: Zunächst denkt man vielleicht, dass es zwei Arten von Objekten gibt: Elemente und Mengen. Nun können aber Mengen Elemente anderer Mengen sein und die Elemente von Mengen können selbst wieder Mengen sein. Als nächstes sieht man eine Hierarchie von Objekten: Elemente, die selbst keine Mengen sind (sogenannte **Urelemente**), Mengen, Mengen von Mengen, ... Aber auch dies ist nicht ausreichend, weil man diese Ebenen auch durchmischen kann, z. B. kann man die zweielementige Menge $\{1, \{2, 3, 5\}\}$ oder die dreielementige Menge $\{\emptyset, \mathbb{N}, \{\mathbb{R}\}\}$ bilden.

Vorsicht: Nicht alles, was man beschreiben oder in Formeln hinschreiben kann, ist eine Menge, da dies schnell zu Widersprüchen führt. Zum Beispiel lässt man die „Menge aller Mengen“ nicht zu, d. h. es gibt keine Menge, die alle Mengen als Elemente hätte (der Grund wird auf Seite 40 erklärt). In der axiomatischen Mengenlehre wird genau festgelegt, welche Beschreibungen Mengen definieren. Die im folgenden Abschnitt beschriebenen Konstruktionen kann man aber sorglos verwenden.

1.2 Elementare Konstruktionen und ihre Eigenschaften

- Zu je zwei Mengen M und N gibt es deren **Schnittmenge** $M \cap N$ von M und N , die folgendermaßen definiert ist: Ein Objekt x ist ein Element von $M \cap N$, wenn es ein Element von M ist *und* ein Element von N ist.

Schnitt $\cap$

Beispiel: $\{2, 3, 5\} \cap \{1, 2, 3, 4\} = \{2, 3\}$.

Falls $M \cap N = \emptyset$, so heißen M und N **disjunkt**.

- Zu je zwei Mengen M und N gibt es deren *Vereinigungsmenge* $M \cup N$ von M und N , die folgendermaßen definiert ist: Ein Objekt x ist ein Element von $M \cup N$, wenn es ein Element von M ist *oder* ein Element von N ist.

Vereinigung $\cup$

Beispiel: $\{2, 3, 5\} \cup \{1, 2, 3, 4\} = \{1, 2, 3, 4, 5\}$.

Die Vereinigung von M und N zu $M \cup N$ heißt *disjunkt*, falls M und N disjunkt sind. Viele schreiben dafür $M \cup N$ statt $M \cup N$, falls es wichtig ist, dass es sich um eine disjunkte Vereinigung handelt.

 $\cup$

- Zu je zwei Mengen M und N gibt es die *Differenzmenge* $M \setminus N$, deren Elemente genau diejenigen Elemente von M sind, die nicht Elemente von N sind. Manche Autoren schreiben auch $M - N$. Die Differenzmenge heißt auch *Komplement(menge) von N in M* . Dies wird insbesondere benutzt, wenn N Teilmenge von M ist. Dann schreibt man auch N^c oder $\complement N$, vorausgesetzt die Menge M , in der das Komplement gebildet wird, ist explizit oder durch den Kontext festgelegt.

Differenz $\setminus$ Komplement $\complement$

Beispiel: $\{2, 3, 5\} \setminus \{1, 2, 3, 4\} = \{5\}$.

Die *symmetrische Differenz* zweier Mengen M und N ist $M \Delta N := (M \setminus N) \cup (N \setminus M)$. Es gilt dann $M \cup N = (M \cap N) \cup (M \Delta N)$.

symmetrische Differenz Δ

Beispiel: $\{2, 3, 5\} \Delta \{1, 2, 3, 4\} = \{1, 4, 5\}$.

- (*Aussonderung*) Wenn M eine Menge ist und E eine mathematisch beschreibbare Eigenschaft, dann gibt es die Teilmenge der Elemente von M , die die Eigenschaft E erfüllen, also die Menge $\{x \in M \mid x \text{ hat die Eigenschaft } E\}$.

- Zu jeder Menge M gibt es die *Potenzmenge* $\mathfrak{P}(M)$ oder $\text{Pot}(M)$, deren Elemente genau die Teilmengen von M sind, d. h. es gilt $X \in \mathfrak{P}(M) \iff X \subseteq M$.

Potenzmenge $\mathfrak{P}(M)$

Beispiel: $\mathfrak{P}(\{1, 2, 3\}) = \{\emptyset, \{1\}, \{2\}, \{3\}, \{1, 2\}, \{1, 3\}, \{2, 3\}, \{1, 2, 3\}\}$

Wenn M endlich ist und n Elemente hat, dann besteht $\mathfrak{P}(M)$ aus 2^n Elementen. Vor allem in der Informatik wird die Potenzmenge von M daher auch mit 2^M bezeichnet.

- Jede endliche Anzahl an Objekten $a_1, \dots, a_n$ kann man zu der Menge $\{a_1, \dots, a_n\}$ zusammenfassen, deren Elemente also gerade die Objekte $a_1, \dots, a_n$ sind.

Wenn die Elemente $a_1, \dots, a_n$ paarweise verschieden sind, wird die Menge $\{a_1, \dots, a_n\}$ manchmal auch *ungeordnetes n -Tupel* genannt.

- Für jede endliche Anzahl an Objekten $a_1, \dots, a_n$ gibt es ein Objekt $(a_1, \dots, a_n)$, das (*geordnete*) *n -Tupel der a_i* , mit folgender Eigenschaft:

 n -Tupel

$$(a_1, \dots, a_n) = (b_1, \dots, b_n) \iff a_1 = b_1, a_2 = b_2, \dots, a_n = b_n$$

Man nennt a_i dann die *i -te Komponente* des Tupels. Also sind zwei geordnete n -Tupel genau dann gleich, wenn sie in jeder Komponente übereinstimmen. Für $n = 2, 3, 4, \dots$ heißen die n -Tupel *Paare*, *Tripel*, *Quadrupel*, *Quintupel*, *Sextupel*, *Septupel* ... Der Name „Tupel“ leitet sich davon ab.

Per Konvention gibt es ein einziges 0-Tupel, nämlich $\emptyset$.

Wenn man weiß, was Abbildungen sind (siehe Seite 14), stellt man sich ein n -Tupel $(a_1, \dots, a_n)$ am besten vor als eine Abbildung von der Menge der Indizes $\{1, \dots, n\}$ in die Menge $\{a_1, \dots, a_n\}$, und zwar die Abbildung, die jedem Index i das Element a_i zuordnet.

Bei einem rein mengentheoretischen Aufbau der Mathematik hat man aber das Problem, dass man geordnete Paare braucht, um Abbildungen definieren zu können. Man muss dafür also eine andere Art von geordnetem Paar finden. Nun beschreibt die oben angegebene Eigenschaft keine eindeutige Auswahl von Objekten, denn man kann durch mengentheoretische Konstruktionen auf verschiedene Weise Mengen mit dieser Eigenschaft finden. Um eine exakte mengentheoretische

Definition der geordneten Tupel zu geben, muss man eine Konstruktionsart auswählen. Welche man wählt und wie man es genau macht, ist aber für Fragen der mathematischen Praxis ohne Belang.

Meistens wählt man das *Kuratowski-Paar*. Hierbei wird (a_1, a_2) als die Menge $\{\{a_1\}, \{a_1, a_2\}\}$ definiert. Man kann nun beweisen, dass diese Konstruktion die Anforderung erfüllt, also dass dann und nur dann $\{\{a_1\}, \{a_1, a_2\}\} = \{\{b_1\}, \{b_1, b_2\}\}$ ist, wenn $a_1 = a_2$ und $b_1 = b_2$. (Dabei muss man beachten, dass der Fall $a_1 = a_2$ auftreten kann.) Tripel und allgemeine n -Tupel kann man dann durch geschachtelte Paare erzeugen, z.B. Tripel als $((a_1, a_2), a_3)$ oder als $(a_1, (a_2, a_3))$. Die naheliegende Verallgemeinerung $\{\{a_1\}, \{a_1, a_2\}, \{a_1, a_2, a_3\}\}$ funktioniert dagegen nicht.

- Für je endlich viele Mengen $M_1, \dots, M_n$ gibt es das mit $M_1 \times \dots \times M_n$ bezeichnete *Mengenprodukt* oder *kartesische Produkt* der Mengen M_i . Dies ist eine Menge, deren Elemente genau die n -Tupel $(m_1, \dots, m_n)$ mit $m_1 \in M_1, \dots, m_n \in M_n$ sind.

Produkt $\times$

Falls $M_1 = \dots = M_n = M$, so schreibt man auch M^n für $M_1 \times \dots \times M_n$ (es sei denn, auf M gibt es auch noch eine Multiplikation, dann wird wegen der Verwechslungsgefahr manchmal eine andere Notation wie z.B. $M^{\times n}$ gewählt).

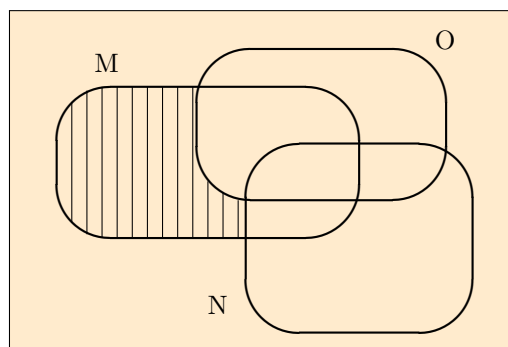
Beispiele: $\{1, 3\} \times \{1, 2, 4\} = \{(1, 1), (1, 2), (1, 4), (3, 1), (3, 2), (3, 4)\}$

$\{a, b\}^3 = \{(a, a, a), (a, a, b), (a, b, a), (a, b, b), (b, a, a), (b, a, b), (b, b, a), (b, b, b)\}$

Wenn $|M_1| = k_1, \dots, |M_n| = k_n$, dann besteht $M_1 \times \dots \times M_n$ aus $k_1 \cdot \dots \cdot k_n$ Elementen. Insbesondere ist $|M|^n$ die Anzahl der Elemente von M^n .

Aus der Konvention über 0-Tupel ergibt sich $M^0 = \{\emptyset\}$ für beliebige Mengen M . Außerdem gilt $\emptyset^n = \emptyset$ für alle $n > 0$. Demgemäß setzt man auch $n^0 = 1$ für alle natürlichen Zahlen n und $0^n = 0$ für alle natürlichen Zahlen $n > 0$.

Für die mengentheoretischen Operationen Schnitt, Vereinigung, Differenz, Produkt gelten Gesetze (Rechenregeln), von denen eine ganze Reihe in Tabelle 1 auf Seite 9 aufgeführt sind. Die meisten dieser Regeln haben die folgende Gestalt: Es treffen zwei Operationen aufeinander und die Regel sagt etwas darüber aus, inwieweit man die Reihenfolge der Operationen vertauschen kann. Zum Beispiel die Regel $(M \setminus N) \setminus O = M \setminus (N \cup O)$, bei der die „innere Operation“ $M \setminus \dots$ auf der linken Seite gewissermaßen nach außen gezogen wird, dabei verändert sich hier aber die andere Operation. Es ist eine gute Übung, die Gültigkeit einiger dieser Regeln nachzuweisen. Man kann dies durch Rechnen mit den Elementen tun (d.h. man überlegt sich: Wann ist etwas ein Element von $(M \setminus N) \setminus O$?, wann ist etwas ein Element von $M \setminus (N \cup O)$?, und zeigt dann, dass beide Bedingungen äquivalent sind). Man kann sich diese Rechnungen anhand von Mengendiagrammen (auch *Venn-Diagramme* genannt) veranschaulichen, siehe Abbildung 1, und sie, wenn man das Prinzip verstanden hat, auch dadurch ersetzen.



Die Elemente einer Menge, z.B. M , denkt man sich innerhalb der mit M bezeichneten geschlossenen Kurve liegend. Die schraffierte Fläche entspricht dann den Elementen von M , die weder in N noch in O liegen. Man kann sich dann bildlich davon überzeugen, dass man sie auf zwei Weisen erhalten kann: Indem man N und O zusammen aus M ausschneidet (was $M \setminus (N \cup O)$ entspricht), oder indem man erst die Menge N aus M herausnimmt und anschließend noch O herausnimmt (was $(M \setminus N) \setminus O$ entspricht.)

Abbildung 1: VENN-DIAGRAMME

Etwas schwieriger zu verstehen sind manche Regeln für Produkte.

Ob Gleichheiten wie $M \times N \times O = (M \times N) \times O$ oder $M \times N \times O = M \times (N \times O)$ gelten oder nicht, hängt von der konkreten Realisierung der geordneten Tupel als Mengen ab. Nach der weit verbreiteten Methode der geschachtelten Kuratowski-Paare gilt nur eine der beiden Gleichungen. Durch „Weglassen von Klammern“ gibt es aber eine natürliche Identifikation von $(M \times N) \times O$ bzw. $M \times (N \times O)$ mit

Tabelle 1: REGELN FÜR MENGENTHEORETISCHE OPERATIONEN

M, N, O, M_i, N_i sind beliebige Mengen:

Kommutativgesetze $M \cap N = N \cap M$
 $M \cup N = N \cup M$

Assoziativgesetze $(M \cap N) \cap O = M \cap (N \cap O)$
 $(M \cup N) \cup O = M \cup (N \cup O)$

Distributivgesetze $(M \cap N) \cup O = (M \cup O) \cap (N \cup O)$
 $(M \cup N) \cap O = (M \cap O) \cup (N \cap O)$

Absorptionsgesetze $(M \cap N) \cup M = M$
 $(M \cup N) \cap M = M$

Rechnen mit $\emptyset$ $M \cup \emptyset = M \setminus \emptyset = M$
 $M \cap \emptyset = M \times \emptyset = \emptyset \times M = \emptyset$
 $M \cup \emptyset = M \setminus \emptyset = M$
 $M \cap \emptyset = M \times \emptyset = \emptyset \times M = \emptyset$

allgemeiner: Falls $n \geq 1$ und $M_i = \emptyset$ für ein $i \in \{1, \dots, n\}$,
so ist $M_1 \times \dots \times M_n = \emptyset$

Regeln von de Morgan $M \setminus (N \cap O) = (M \setminus N) \cup (M \setminus O)$
 $M \setminus (N \cup O) = (M \setminus N) \cap (M \setminus O)$

alternativ: falls $N, O \subseteq M$ und das Komplement in M gebildet wird:
 $(N \cap O)^c = N^c \cup O^c$
 $(N \cup O)^c = N^c \cap O^c$

Regeln für Differenzen $M \setminus N = M \setminus (M \cap N)$
 $M \setminus (M \setminus N) = M \cap N$
 $(M \setminus N) \setminus O = M \setminus (N \cup O)$

Regeln für Teilmengen $M \subseteq N \iff M \cap N = M \iff M \cup N = N$
 $M \subseteq N \subseteq O \implies M \subseteq O$

Distributivgesetze $(M \cap N) \times O = (M \times O) \cap (N \times O)$
 $(M \cup N) \times O = (M \times O) \cup (N \times O)$
 $(M \setminus N) \times O = (M \times O) \setminus (N \times O)$
 $M \times (N \cap O) = (M \times N) \cap (M \times O)$
 $M \times (N \cup O) = (M \times N) \cup (M \times O)$
 $M \times (N \setminus O) = (M \times N) \setminus (M \times O)$
 $(M_1 \times N_1) \cap (M_2 \times N_2) = (M_1 \cap M_2) \times (N_1 \cap N_2)$

$M \times N \times O$, indem man $((m, n), o)$ bzw. $(m, (n, o))$ mit (m, n, o) gleichsetzt. Allgemeiner kann man $((m_1, \dots, m_l), (m_{l+1}, \dots, m_{l+k}))$ mit $(m_1, \dots, m_l, m_{l+1}, \dots, m_{l+k})$ gleichsetzen und damit für praktische Belange $(M_1 \times \dots \times M_l) \times (M_{l+1} \times \dots \times M_{l+k})$ mit $M_1 \times \dots \times M_{l+k}$ identifizieren. Im Spezialfall $M_1 = \dots = M_{l+k} = M$ kann man also $M^l \times M^k$ mit M^{l+k} identifizieren.

Unintuitiv ist das Verhalten von „leeren Produkten“: Wegen der Festsetzung, dass $\emptyset$ das einzige 0-Tupel ist, gilt $M^0 = \{\emptyset\}$ für jede Menge M , und damit auch $\emptyset^0 = \{\emptyset\}$. Allgemeiner ist $M_1 \times \dots \times M_0 = \{\emptyset\}$ (der Index läuft hier von 1 bis 0, es ist also ein „leeres Produkt“!). Man kann dies für den Anfang als eine Konvention verstehen.

1.3 Relationen

Eine (*n-stellige*) *Relation* R zwischen Mengen $M_1, \dots, M_n$ ist von der Grundidee her etwas, von dem man sagen kann, dass es auf Elemente $m_1 \in M_1, \dots, m_n \in M_n$ zutrifft oder nicht. Eine zweistellige Relation außerhalb der Mathematik ist etwa die *Verheiratet-Sein-Relation* zwischen der Menge der Männer und der Menge der Frauen¹; ein Beispiel einer dreistelligen Relation ist: „... ist Kind von ... und ...“. Mathematische (zweistellige) Relationen sind beispielsweise die *Kleiner-oder-gleich-Relation* $\leq$ zwischen natürlichen Zahlen oder die *Teilmengenrelation* $\subseteq$ zwischen zwei Mengen von Mengen oder die *Elementrelation* $\in$ zwischen einer Menge M und ihrer Potenzmenge $\mathfrak{P}(M)$. Ein mathematisches Beispiel einer höherstelligen Relation ist die $(n+1)$ -stellige Relation „... ist das geordnete n -Tupel von ...“ (und für die zweiten Pünktchen müssen n Objekte eingesetzt werden).

Man kann eine Relation R mengentheoretisch mit einer Teilmenge von $M_1 \times \dots \times M_n$ gleichsetzen, nämlich mit der Menge der Tupel, auf die die Relation zutrifft, wobei die Mengen $M_1, \dots, M_n$ zur Definition hinzugehören. Die übliche mathematische Definition ist daher, dass eine n -stellige Relation R ein $(n+1)$ -Tupel $(G, M_1, \dots, M_n)$ von Mengen ist mit der Zusatzbedingung $G \subseteq M_1 \times \dots \times M_n$. Die Menge G wird dabei der *Graph der Relation* R genannt und meist mit G_R oder Γ_R bezeichnet.

Graph G_R

- Oft werden bei der Bezeichnung einer Relation die Mengen $M_1, \dots, M_n$ weggelassen und nur der Graph angegeben.
- Oft unterscheidet man nicht sauber zwischen einer Relation und ihrem Graphen und schreibt „ $R \subseteq M_1 \times \dots \times M_n$ sei eine Relation“ und bezeichnet dann mit R sowohl die Relation als auch ihren Graphen (und das werde ich der Einfachheit halber später auch tun).

Wenn $M_1 \subseteq N_1, \dots, M_n \subseteq N_n$, dann ist eine Teilmenge G von $M_1 \times \dots \times M_n$ auch Teilmenge von $N_1 \times \dots \times N_n$. Also ist G sowohl Graph einer Relation R_M zwischen $M_1, \dots, M_n$ als auch einer Relation R_N zwischen $N_1, \dots, N_n$. Wenn die Mengen $M_1, \dots, M_n$ nun verschieden von den Mengen $N_1, \dots, N_n$ sind, sind R_M und R_N verschiedene Relationen. R_M heißt dann *die auf $M_1, \dots, M_n$ eingeschränkte Relation R_N* .

Falls $M_1 = \dots = M_n = M$, spricht man auch von einer *n-stelligen Relation auf M* . Wenn ein Tupel $(m_1, \dots, m_n)$ im Graphen einer Relation R liegt, sagt man: „ R trifft auf $(m_1, \dots, m_n)$ zu“ oder „ $m_1, \dots, m_n$ stehen in der Relation R zueinander“ oder „ $(m_1, \dots, m_n)$ liegt in R “, oder Ähnliches, und schreibt dafür $Rm_1 \dots m_n$ oder $R(m_1, \dots, m_n)$ oder $(m_1, \dots, m_n) \in R$.

Eine einstellige Relation auf einer Menge M ist eine Eigenschaft, die Elemente von M haben oder nicht. Zum Beispiel ist „Primzahl sein“ eine einstellige Relation auf den natürlichen Zahlen. Der Graph einer einstelligen Relation auf M ist nichts anderes als eine Teilmenge von M .

Zweistellige Relationen

Von besonderer Bedeutung sind die zweistelligen oder *binären Relationen*. Für eine binäre Relation $R \subseteq M \times N$ heißt M der *Definitions-* oder *Vorbereich* und N der *Werte-* oder *Nachbereich*. Das Vertauschen der beiden Rollen führt zu der *Umkehrrelation* R^{-1} zwischen N und M , deren Graph gerade der „gespiegelte“ Graph von R ist, d. h. $(n, m) \in R^{-1} \iff (m, n) \in R$. Beispielsweise ist die „größer“-Relation auf den natürlichen Zahlen die Umkehrrelation der „kleiner“-Relation.

Definitions-,
Wertebereich
 R^{-1}

Bei binären Relationen schreibt man gerne mRn dafür, dass R auf (m, n) zutrifft, insbesondere wenn für R ein Symbol wie z. B. $\leq$ oder $\sim$ verwendet wird.

Sei nun R eine binäre Relation auf M (Definitions- und Wertebereich stimmen also überein). Es folgt nun in drei Gruppen eine Reihe von Eigenschaften binärer Relationen, die in der Mathematik immer wieder betrachtet werden.

¹2010 war der Begriff „verheiratet“ noch für Ehen zwischen Männern und Frauen reserviert.

Man kann eine binäre Relation so visualisieren, wie man es von der Schule her mit Funktionen gewöhnt ist: Man trägt auf einer horizontalen Achse die Elemente von M und auf einer vertikalen Achse die Elemente von N ab, und zeichnet Punkte an den Stellen (m, n) ein, die im Graph der Relation liegen.

Im Beispiel rechts: der Graph der Teilbarkeitsrelation eingeschränkt auf $\{1, 2, 3, 4, 5, 6\}$.

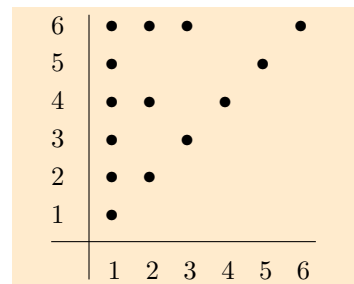


Abbildung 2: Visualisierung von binären Relationen

In der ersten Gruppe geht es darum, ob Elemente zu sich selbst in Relation stehen.

- R heißt **reflexiv**, wenn für alle $m \in M$ gilt, dass R auf (m, m) zutrifft.
- R heißt **irreflexiv** oder **antireflexiv**, wenn es kein $m \in M$ gibt, so dass R auf (m, m) zutrifft. („Irreflexiv“ ist also etwas anderes als „nicht reflexiv“!)

reflexiv

Die Alltagsrelation (zwischen Menschen) „den gleichen Namen tragen wie“ ist reflexiv, die Relation „verheiratet sein mit“ dagegen ist irreflexiv, da niemand mit sich selbst verheiratet ist.

Ein typisches mathematisches Beispiel einer reflexiven Relation ist die „kleiner oder gleich“-Relation auf den natürlichen Zahlen, wohingegen die „(echt) kleiner“-Relation auf den natürlichen Zahlen irreflexiv ist. Eine Relation kann sehr wohl weder reflexiv noch irreflexiv sein, zum Beispiel die Relation „disjunkt sein zu sich selbst“ auf der Potenzmenge einer nicht-leeren Menge, etwa $\mathfrak{P}(\{a\})$. Die leere Menge ist wegen $\emptyset \cap \emptyset = \emptyset$ disjunkt zu sich selbst, wohingegen $\{a\}$ wegen $\{a\} \cap \{a\} \neq \emptyset$ nicht disjunkt zu sich selbst ist.

In der zweiten Gruppe von Eigenschaften geht es darum, ob man die Relation „herumdrehen“ kann, also ob man von der Relation auf die Umkehrrelation schließen kann.

- R heißt **symmetrisch**, wenn für alle $m, n \in M$ gilt, dass R genau dann auf (m, n) zutrifft, wenn es auf (n, m) zutrifft.
- R heißt **asymmetrisch**, wenn es keine Elemente $m, n \in M$ gibt, so dass R auf (m, n) und auf (n, m) zutrifft.
- R heißt **antisymmetrisch**, wenn es keine verschiedenen Elemente $m, n \in M$ gibt, so dass R auf (m, n) und auf (n, m) zutrifft.

symmetrisch

Zunächst wieder Alltagsrelationen zwischen Mengen: Die *Verheirate-Sein*-Relation ist symmetrisch, dagegen ist die *Verliebt-Sein*-Relation nicht symmetrisch, aber zum Glück auch nicht asymmetrisch oder antisymmetrisch.

Mathematische Beispiele auf einer Potenzmenge $\mathfrak{P}(M)$: Die Relation „hat genauso viele Element wie“ ist symmetrisch; die Relation „ist Teilmenge von“ ist antisymmetrisch und die Relation „ist echte Teilmenge von“ ist asymmetrisch (Es kann sein, dass $A \subseteq B$ und $B \subseteq A$, nämlich genau dann, wenn $A = B$, aber es kommt nie vor, dass $A \subset B$ und $B \subset A$).

Schließlich noch eine wichtige Eigenschaft, bei der es um drei Elemente geht:

- R heißt **transitiv**, wenn für alle $m_1, m_2, m_3 \in M$ gilt: Wenn R auf (m_1, m_2) und auf (m_2, m_3) zutrifft, dann auch auf (m_1, m_3) .

transitiv

Eine transitive Alltagsrelation ist „schwerer sein als“; nicht transitiv ist die Freundschafts-Relation: Ein Freund eines Freundes ist nicht immer selbst ein Freund.

Mathematische Beispiele sind die Ordnungsrelationen „*kleiner*“ und „*kleiner oder gleich*“ auf den natürlichen Zahlen. Eine nicht transitive Relation ist z.B. „*ein gemeinsames Element haben*“ auf $\mathfrak{P}(\{a, b, c, d\})$. Die Mengen $\{a, b\}$ und $\{b, c\}$ haben b als gemeinsames Element, die Mengen $\{b, c\}$ und $\{c, d\}$ haben c als gemeinsames Element, aber die Mengen $\{a, b\}$ und $\{c, d\}$ haben kein gemeinsames Element.

Nun werden einige wichtige Arten von binären Relationen auf einer Menge M beschrieben:

Ordnungsrelationen

Eine *totale* oder *lineare Ordnungsrelation* auf M ist eine reflexive, antisymmetrische, transitive binäre Relation auf M , die zudem *total* ist: Dies bedeutet, dass für alle $m, n \in M$ die Relation R auf (m, n) zutrifft oder auf (n, m) . Beispiel ist die wohlbekannte Ordnungsrelation $\leq$ auf den natürlichen Zahlen, d. h. die Relation, die auf m und n dann und nur dann zutrifft, wenn $m \leq n$.

Statt „Ordnungsrelation“ sagt man auch oft kurz „Ordnung“.

Es gibt nun mehrere Abschwächungen von totalen Ordnungen. Eine *Halbordnung* oder *partielle Ordnung* auf M ist eine reflexive, antisymmetrische und transitive binäre Relation auf M . Ein Beispiel ist die Teilmengenbeziehung $\subseteq$ auf der Potenzmenge einer Menge. (In einer partiellen Ordnung $\leq$ können Elemente m und n also *unvergleichbar* sein, d. h. es gilt weder $m \leq n$ noch $n \leq m$.)

Quasi- oder *Präordnungen* sind reflexive und transitive binäre Relationen. (In einer Präordnung $\leq$ folgt also aus $m \leq n$ und $n \leq m$ nicht unbedingt $m = n$.)

Ob das Wort „Ordnung“ ohne Zusatz für eine partielle oder eine totale Ordnung steht, hängt vom Autor ab.

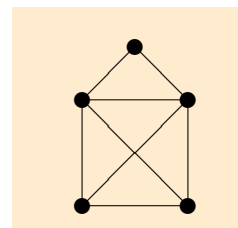
Zu jeder totalen oder partiellen Ordnungsrelation $\leq$ gibt es die zugehörige *strikte Ordnungsrelation*, z. B. auf den natürlichen Zahlen die mit $<$ bezeichnete Relation „*echt kleiner als*“. Sie entsteht dadurch, dass man die Ordnung $\leq$ irreflexiv macht, d. h. $<$ stimmt für Paare (m, n) mit $m \neq n$ mit der Ordnungsrelation $\leq$ überein, trifft aber auf Paare (m, m) nie zu. Eine partielle strikte Ordnung ist eine irreflexive, asymmetrische und transitive binäre Relation. Beispiel hierfür ist die „*echte Teilmenge*“-Relation $\subset$ auf der Potenzmenge einer Menge. Eine totale strikte Ordnung ist damit eine irreflexive, asymmetrische und transitive binäre Relation, die zudem für verschiedene m und n entweder auf (m, n) oder auf (n, m) zutrifft.

Achtung: Entgegen dem normalen Sprachgebrauch ist eine „strikte Ordnung“ keine „Ordnung“! (Normalerweise versucht man solche sprachlichen Phänomene in der Mathematik zu vermeiden.)

Man schreibt für Ordnungsrelationen gerne das Symbol $\leq$ oder ähnliche Symbole wie $\subseteq$, $\preceq$, $\sqsubseteq$. Für die Umkehrrelation wird dann das gespiegelte Symbol benutzt, also $\geq$, $\supseteq$, $\succcurlyeq$, $\sqsupseteq$. Für die zugehörige strikte Relation benutzt man oft das Symbol ohne den Gleichheitsstrich, also $<$, $\subset$, $\prec$, $\subsetneq$. Allerdings werden auch immer wieder für nicht-strikte Ordnungsrelationen solche Symbole ohne Gleichheitsstrich verwendet, etwa in Version 2 der Teilmengebezeichnungen.

Graphen

In einer anderen Verwendungsweise als bisher wird mit *Graph* auch eine irreflexive, symmetrische binäre Relation auf M bezeichnet. Solche Graphen kann man aufmalen, indem man die Elemente von M als dicke Punkte malt, und eine Verbindungslinie zwischen ihnen zieht, wenn sie in Relation zueinander stehen.



Beliebige binäre Relationen auf einer Menge M heißen auch *gerichtete Graphen*. Hier werden die Verbindungslinien in der Veranschaulichung als Pfeile von m nach n gemalt, falls die Relation auf (m, n) zutrifft.

Äquivalenzrelationen

Eine **Äquivalenzrelation** auf M ist eine reflexive, symmetrische und transitive binäre Relation auf M . Für Äquivalenzrelationen benutzt man gerne symmetrische Symbole, die dem Gleichheitszeichen mehr oder weniger ähneln, wie $\equiv$, $\approx$, $\sim$, $\dots$.

Ist $\sim$ eine Äquivalenzrelation auf M , so heißt die Menge der Elemente von M , die zu einem gegebenen $m \in M$ in Relation stehen, die **Äquivalenzklasse** von m . Für die Äquivalenzklassen gibt es viele verschiedene Notationen; ich schreibe hier $m/\sim$ für die Äquivalenzklasse von m . Das Element m heißt **Repräsentant** seiner Klasse. Zwei Äquivalenzklassen sind entweder gleich oder disjunkt; die Äquivalenzklassen bilden also eine Familie $(m/\sim)_{m \in M}$ von nicht-leeren Teilmengen von M , die paarweise entweder gleich oder disjunkt sind (zum Begriff der Familie siehe Seite 17). Solch eine Familie wird **Partition** von M genannt. Ist umgekehrt eine Partition $(M_i)_{i \in I}$ von M gegeben, so definiert $x \in M_i \iff y \in M_i$ eine Äquivalenzrelation $x \sim y$ auf M , deren Klassen gerade die Teilmengen M_i der Partition sind.

Ein Alltagsbeispiel: Betrachtete wird die Menge der Schüler einer Schule und darauf die binäre Relation „in der gleiche Klasse sein“. Diese Relation ist eine Äquivalenzrelation; die Äquivalenzklassen sind genau die Schulklassen. Jeder Schüler ist Repräsentant seiner Klasse.

Die Bildung von Äquivalenzklassen ist entscheidend für die mathematische Begriffsbildung und den in der Mathematik häufigen Abstraktionsprozess. An einem nicht-mathematischen Beispiel erläutert: Wenn man eine Menge von farbigen Gegenständen hat, kann man darauf die binäre Relation der Gleichfarbigkeit betrachten. Dies ist eine Äquivalenzrelation; die Äquivalenzklasse eines Gegenstands besteht aus den dazu gleichfarbigen Gegenständen. Diese Äquivalenzklasse steht damit gewissermaßen für die Farbe: Farbe ist das, was die Gegenstände der Äquivalenzklasse gemeinsam haben. Auf diese Weise kommt man von der konkreten Relation der Gleichfarbigkeit (die man bestimmen kann auch ohne Farben oder Namen von Farben zu kennen) zu den abstrakten Farben.

So würde man auch vorgehen, wenn man die Tiere in einer unbekannten Welt klassifizieren wollte: Man sortiert sie nach der Äquivalenzrelation „gleiche Art“ (wozu man die Arten nicht zu kennen braucht) und gibt dann erst den Äquivalenzklassen, also den Arten, die Namen.

Ein bekanntes mathematisches Beispiel ist die Konstruktion der rationalen Zahlen aus den ganzen Zahlen: Verschieden aussehende Brüche wie $\frac{2}{3}$ und $\frac{4}{6}$ können für die gleiche rationale Zahl stehen; ein Bruch $\frac{a}{b}$ ist damit eigentlich eine Äquivalenzklasse von Paaren ganzer Zahlen (a, b) unter der Äquivalenzrelation $(a, b) \sim (c, d) \iff ad = bc$ (wobei man noch einiges dazusagen muss, da man durch 0 nicht teilen kann).

Wenn $\sim$ eine Äquivalenzrelation auf der Menge M ist, kann man auf der Menge der Äquivalenzklassen Eigenschaften und Relationen definieren, indem man Eigenschaften und Relationen der Repräsentanten überträgt. Wenn also eine Eigenschaft E gegeben ist, die Elemente von M haben können oder nicht, dann definiert man die Eigenschaft E' für die Äquivalenzklassen durch die folgende Festlegung:

Die Äquivalenzklasse $x/\sim$ habe die Eigenschaft E' , wenn x die Eigenschaft E hat.

Nun hat aber eine Äquivalenzklasse i. Allg. viele Repräsentanten, also ist diese Definition nur dann sinnvoll, wenn sie nicht von der Auswahl des Repräsentanten abhängt, d. h. wenn auch alle anderen Elemente y aus der Äquivalenzklasse $x/\sim$ die Eigenschaft E haben, sobald x die Eigenschaft E besitzt. In diesem Fall sagt man, dass die Definition **unabhängig von der Wahl der Repräsentanten** oder kurz **wohldefiniert** ist.

Ein Beispiel: Man definiert, dass ein Bruch $\frac{a}{b}$ (ungleich 0) positiv ist, wenn die beiden Komponenten des Repräsentanten (a, b) gleiches Vorzeichen haben. Dies ist wohldefiniert.

Die versuchsweise Definition, einen Bruch $\frac{a}{b}$ gerade zu nennen, wenn a gerade ist, scheitert dagegen an der mangelnden Wohldefiniertheit.

Natürlich gibt es auch Definitionen von Eigenschaften von Äquivalenzklassen, die anderer Art sind. Zum Beispiel kann man definieren, dass eine Äquivalenzklasse $x/\sim$ eine Eigenschaft E habe, wenn es einen Repräsentanten y gibt, der eine Eigenschaft E' hat. Dann muss keine Wohldefiniertheit nachgewiesen werden. Der Unterschied liegt hier in der Art der Quantifizierung: Einmal verlangt man, dass *alle* Repräsentanten die Eigenschaft E haben, das andere Mal, dass es *mindestens einen* solch einen Repräsentanten gibt.

Als Übung im Umgang mit den Begriffen kann man Folgendes nachprüfen: Wenn $\leq$ eine Präordnung auf M ist, so definiert „ $x \leq y$ und $y \leq x$ “ eine Äquivalenzrelation $x \sim y$. Durch die Festlegung $x/\sim \leq' y/\sim$, falls $x \leq y$, ist dann auf der Menge der Äquivalenzklassen eine wohldefinierte partielle Ordnung $\leq'$ gegeben.

Eine besondere Äquivalenzrelation ist die Identität, also die Relation „mit etwas identisch sein“, die mit dem Gleichheitszeichen $=$ geschrieben wird. Ihre Äquivalenzklassen sind also alle ein-elementig.

Die Identität hat offensichtlich folgende besondere Eigenschaft: Wenn $m_1 = n_1, \dots, m_k = n_k$ und R eine k -stellige Relation, dann trifft R genau dann auf $(m_1, \dots, m_k)$ zu, wenn sie auf $(n_1, \dots, n_k)$ zutrifft. Eine Äquivalenzrelation $\sim$, die bezüglich einer Menge $\mathcal{R}$ von Relationen diese Eigenschaft hat, heißt *Kongruenzrelation* bezüglich $\mathcal{R}$. In diesem Fall sind alle Relationen in $\mathcal{R}$ repräsentantenunabhängig und können auf die Menge der $\sim$ -Äquivalenzklassen übertragen werden. In solchen Situationen „identifiziert“ man gerne die Objekte einer $\sim$ -Äquivalenzklasse. Dies kann man entweder so verstehen, dass man „gleich“ statt „äquivalent“ sagt und $=$ statt $\sim$ schreibt. Man kann es aber auch so verstehen, dass man mit den Äquivalenzklassen von $\sim$ statt mit den ursprünglichen Objekten arbeitet, aber so über die Äquivalenzklassen spricht, als wären es noch die ursprünglichen Objekte.

1.4 Abbildungen

Eine *Abbildung* von einer Menge M in eine Menge N ist von der Grundidee her eine Zuordnung, die jedem Element von M genau ein Element aus N zuordnet. Solch eine Zuordnung definiert eine Relation, und zwar trifft die Relation genau dann auf $m \in M$ und $n \in N$ zu, wenn n das Element ist, das m zugeordnet wird. Eine Abbildung von M nach N ist daher definiert als eine *linkstotale* und *rechtseindeutige* Relation zwischen M und N . Dabei bedeutet *linkstotal*, dass es zu jedem $m \in M$ (mindestens) ein $n \in N$ gibt, das mit m in der Relation steht. *Rechtseindeutig* bedeutet, dass es zu jedem $m \in M$ höchstens ein $n \in N$ gibt, das mit n in Relation steht. Zusammen gibt es also bei einer Abbildung zu jedem $m \in M$ genau ein $n \in N$, das zu m in Relation steht. Dieses Element n heißt dann *das Bild von m unter f* und man schreibt $n = f(m)$. Umgekehrt ist m *ein Urbild* von n . Ein Element von M kann also nur ein Bild, ein Element von N aber mehrere Urbilder haben. Es kann aber auch sein, dass Elemente von N gar keine Urbilder haben.

Man schreibt $f : M \rightarrow N$ dafür, dass f eine Abbildung von M nach N ist. Manche Abbildungen sind durch *Abbildungsvorschriften* gegeben, das sind (uniforme) Regeln, die angeben, wie man von jedem m zu seinem Bild $f(m)$ kommt. Solche eine Vorschrift könnte z.B. „multipliziere mit 2“ (auf den natürlichen Zahlen) sein; sie wird dann $f : \mathbb{N} \rightarrow \mathbb{N}, n \mapsto 2n$ geschrieben. Es versteht sich aus dieser Schreibweise, dass n eine Variable für Elemente von $\mathbb{N}$ ist, und dass die Funktion f jedes $n \in \mathbb{N}$ auf $2n$ abbildet. Achtung: Da eine Abbildung eine Relation ist, gehören Definitions- und Wertebereich zur Abbildung hinzu. Eine Abbildungsvorschrift allein definiert noch keine Abbildung (manchmal ist man aber nachlässig, wenn sich Definitions- und Wertebereich aus dem Kontext ergeben).

Für viele Mathematiker sind die Begriffe *Abbildung* und *Funktion* gleichbedeutend. Bisweilen wird der Begriff „Funktion“ aber eingeschränkter benutzt, z.B. nur für solche Abbildungen, deren Wertebereich die reellen oder komplexen Zahlen sind.

Wegen der Linkstotalität ist der Definitionsbereich einer Abbildung durch ihren Graphen festgelegt. Die Vorschrift $x \mapsto \frac{1}{x}$ definiert also keine Abbildung auf den reellen Zahlen, sondern nur

von $\mathbb{R} \setminus \{0\}$ nach $\mathbb{R}$. Den Wertebereich kann man dagegen nicht aus dem Graphen bestimmen. Dennoch wird auch eine Abbildung oft mit ihrem Graphen identifiziert, wenn der Wertebereich aus dem Kontext klar wird oder keine Rolle spielt. Die im Folgenden definierten Eigenschaften „surjektiv“ und „bijektiv“ hängen aber vom genauen Wertebereich ab.

Surjektiv, injektiv, bijektiv

- Eine Abbildung $f : M \rightarrow N$ heißt *surjektiv*, wenn sie *rechtstotal* ist, d. h. wenn es zu jedem $n \in N$ (mindestens) ein $m \in M$ gibt, das auf n abgebildet wird.
- Eine Abbildung $f : M \rightarrow N$ heißt *injektiv*, wenn sie *linkseindeutig* ist, d. h. wenn es zu jedem $n \in N$ höchstens ein $m \in M$ gibt, das auf n abgebildet wird.
- Eine Abbildung $f : M \rightarrow N$ heißt *bijektiv*, wenn sie injektiv und surjektiv ist.

Die Abbildung $\mathbb{Z} \rightarrow \mathbb{N}, x \mapsto |x|$ ist surjektiv, aber nicht injektiv.

Die Abbildung $\mathbb{Z} \rightarrow \mathbb{Z}, x \mapsto |x|$ ist weder surjektiv, noch injektiv.

Die Abbildung $\mathbb{N} \rightarrow \mathbb{N}, x \mapsto 2x$ ist injektiv, aber nicht surjektiv.

Die Abbildung $\mathbb{Q} \rightarrow \mathbb{Q}, x \mapsto 2x$ ist bijektiv.

Eingeschränkte Abbildungen und Bilder

Wenn $f : M \rightarrow N$ eine Abbildung ist, dann heißt die Menge

$$\text{im}(f) := \{n \in N \mid \text{es gibt ein } m \in M \text{ mit } f(m) = n\}$$

das *Bild* (engl: *image*) der Abbildung. Wenn man eine Abbildung auf ihr Bild einschränkt, also die auf $M \times \text{im}(f)$ eingeschränkte Relation betrachtet, wird die Abbildung surjektiv. Man kann eine Abbildung auch in ihrem Definitionsbereich einschränken, d. h. man betrachtet für $M' \subseteq M$ die auf $M' \times N$ eingeschränkte Relation. Diese heißt dann *die auf M' eingeschränkte Abbildung* $f|_{M'}$ oder $f|_M$.

$f|_M$

Zu jeder Abbildung $f : M \rightarrow N$ gibt es eine zugehörige Abbildung (so etwas nennt man dann gerne eine „von f induzierte Abbildung“)

$$\mathfrak{P}(M) \rightarrow \mathfrak{P}(N), X \mapsto \text{im}(f|_X) = \{n \in N \mid \text{es gibt ein } m \in X \text{ mit } f(m) = n\}$$

Diese Abbildung wird der Einfachheit halber meistens auch mit f bezeichnet wird, d. h. man schreibt $f(X)$ für $\text{im}(f|_X)$. Man nennt $f(X)$ dann auch *das Bild von X unter f* . Es gibt folgenden Zusammenhang: Für $X \subseteq M$ ist $f(X) = \{f(x) \mid x \in X\}$. Insbesondere ist $f(\{x\}) = \{f(x)\}$.

Bei merkwürdigen Mengen kann es vorkommen, dass ein Element auch eine Teilmenge ist, z. B. ist $\{2\}$ sowohl ein Element als auch eine Teilmenge der Menge $M = \{2, \{2\}\}$. In diesem Fall ist für eine auf M definierte Abbildung f die Notation $f(\{2\})$ zweideutig. Mengentheoretiker schreiben daher $f[X]$ für das Bild von X unter f .

Umkehrabbildungen und Urbilder

Wenn eine Abbildung $f : M \rightarrow N$ bijektiv ist, dann gibt es auch zu jedem $n \in N$ genau ein $m \in M$ mit $f(m) = n$. Dies bedeutet, dass die Umkehrrelation von f (Erinnerung: f ist auch eine Relation!) auch eine Abbildung ist, die mit f^{-1} bezeichnete *Umkehrabbildung* von f . Es gilt dann also: $f(m) = n \iff f^{-1}(n) = m$. Die ursprüngliche Abbildung ist dann wieder die Umkehrabbildung ihrer Umkehrabbildung, d. h. es gilt $(f^{-1})^{-1} = f$.

f^{-1}

Wenn f nur injektiv ist, kann man f^{-1} auch für die Umkehrabbildung der auf ihr Bild eingeschränkten Abbildung f schreiben, also $f^{-1} : \text{im}(f) \rightarrow M$.

Ähnlich wie man bei einer Abbildung $f : M \rightarrow N$ den gleichen Buchstaben auch für die induzierte Abbildung $\mathfrak{P}(M) \rightarrow \mathfrak{P}(N)$ benutzt, schreibt man auch f^{-1} für eine gewisse induzierte Abbildung. Wenn nämlich $f : M \rightarrow N$ eine beliebige (d. h. nicht notwendig injektive) Abbildung ist, definiert man

$$f^{-1} : \mathfrak{P}(N) \rightarrow \mathfrak{P}(M), Y \mapsto \{m \in M \mid f(m) \in Y\}$$

Das heißt, $f^{-1}(Y)$ ist die größte Teilmenge von M , deren Bild in Y liegt. Man nennt sie *das Urbild von Y unter f* . (Man beachte den Unterschied: Ein Element von N kann mehrere Urbilder haben; eine Teilmenge von N hat nur ein Urbild).

Sauberer wäre wieder die Schreibweise $f^{-1}[Y]$ für das Urbild von Y unter f .

Wenn $f : M \rightarrow N$ injektiv ist, steht f^{-1} nun für zwei Sachen:

- zum einen für die Umkehrabbildung $f^{-1} : \text{im}(f) \rightarrow N$,
- zum anderen für die gerade auf $\mathfrak{P}(N)$ definierte Abbildung.

Es gibt dann wieder einen Zusammenhang: Für $Y \subseteq N$ ist $f^{-1}(Y) = \{f^{-1}(y) \mid y \in Y\}$. Insbesondere: Falls $y = f(x)$ für ein $x \in M$, ist $f^{-1}(y) = x$ und $f^{-1}(\{y\}) = \{x\}$.

Die Schreibweisen und Konventionen um f^{-1} sind alle mit großer Vorsicht zu genießen, weil es leicht zu Missverständnissen kommen kann.

Wenn $f : M \rightarrow N$ injektiv ist, dann ist jetzt auf zwei Weisen eine Abbildung $f^{-1} : \mathfrak{P}(\text{im}(f)) \rightarrow \mathfrak{P}(M)$ definiert worden: einmal als die von der Umkehrabbildung f^{-1} induzierte Abbildung und einmal als die „Urbild-Abbildung“. Zum Glück stimmen beide Definitionen überein, geben also dieselbe Abbildung.

Wenn f nicht injektiv ist, besteht $f^{-1}(\{y\})$ in der Regel aus mehr als einem Element. Da in dieser Situation die Umkehrabbildung nicht definiert ist, schreibt man aus Bequemlichkeit gerne $f^{-1}(y)$ statt $f^{-1}(\{y\})$.

Sei $f : \mathbb{N} \rightarrow \mathbb{N}$ die injektive Abbildung $x \mapsto 2x$. Dann gibt es die Umkehrabbildung $f^{-1} : \text{im}(f) = \{n \in \mathbb{N} \mid n \text{ gerade}\} \rightarrow \mathbb{N}$, $n \mapsto \frac{n}{2}$. Es ist also $f^{-1}(4) = 2$, wohingegen $f^{-1}(5)$ nicht definiert ist, da 5 nicht im Bild von f liegt. Dagegen ist $f^{-1}(\{4\}) = \{2\}$ und $f^{-1}(\{5\}) = \emptyset$.

Weiter ist $f(\{2, 3, 4, 5\}) = \{4, 6, 8, 10\}$, also $f^{-1}(\{4, 6, 8, 10\}) = \{2, 3, 4, 5\}$, da f injektiv ist. Andererseits ist $f^{-1}(\{2, 3, 4, 5\}) = \{1, 2\}$.

Sei $g : \mathbb{Z} \rightarrow \mathbb{Z}$ die Abbildung $x \mapsto |x|$. Dann ist $g^{-1}(\{5\}) = \{5, -5\}$ und $g^{-1}(\{3, 4, 5\}) = \{-5, -4, -3, 3, 4, 5\}$ und $g^{-1}(\{-3, -4, -5\}) = \emptyset$.

Eigentlich ist $g^{-1}(5)$ nicht definiert, wird aber manchmal für $g^{-1}(\{5\})$ geschrieben.

Verknüpfungen von Abbildungen

Wenn zwei Abbildungen $f : M \rightarrow N$ und $g : N \rightarrow O$ gegeben sind, dann kann man diese „hintereinander ausführen“, indem man ein Element $m \in M$ erst mit f auf $f(m) \in N$ abbildet und dann mit g auf $g(f(m)) \in O$. Tatsächlich ist die Menge $\{(m, g(f(m))) \mid m \in M\}$ der Graph einer Abbildung $M \rightarrow O$, die *Verknüpfung* oder *Komposition* der beiden Abbildungen f und g genannt wird und mit $g \circ f$ bezeichnet wird.

$g \circ f$

Achtung auf die Reihenfolge: Bei $g \circ f$ wird als Abbildung zunächst f und dann g ausgeführt, d. h. die Reihenfolge dreht sich um. Es ist aber $(g \circ f)(m) = g(f(m))$, hier stimmt die Reihenfolge also.)

Die Komposition von Abbildungen ist assoziativ, d. h. es gilt

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

bei geeigneten Abbildungen, also für Abbildungen $f : M \rightarrow N, g : N \rightarrow O, h : O \rightarrow P$.

Die Verknüpfung zweier injektiver (surjektiver/bijektiver) Abbildungen ist wieder injektiv (surjektiv/bijektiv). Wenn $g \circ f$ injektiv ist, dann ist auch f injektiv; wenn $g \circ f$ surjektiv ist, dann ist auch g surjektiv.

Wenn $f : M \rightarrow N$ und $g : N \rightarrow O$ bijektiv sind und also auch $g \circ f$, dann gilt

$$(g \circ f)^{-1} = f^{-1} \circ g^{-1} : O \rightarrow M.$$

Für jede Menge M gibt es als besondere Abbildung die *identische Abbildung* oder *Identität* $\text{id}_M : M \rightarrow M, m \mapsto m$. Wenn $f : M \rightarrow N$ bijektiv ist, dann gilt $f^{-1} \circ f = \text{id}_M$ und $f \circ f^{-1} = \text{id}_N$. Es gilt auch eine Umkehrung:

- Eine Abbildung $g : N \rightarrow O$ ist injektiv
 $\iff$ es gibt ein *Linksinverses* $h : \text{im}(g) \rightarrow M$ mit $h \circ g = \text{id}_N$ (nämlich $h = g^{-1}$)
 $\iff$ g ist *rechtskürzbar*, d. h. wenn $f_1, f_2 : M \rightarrow N$ Abbildungen sind mit $g \circ f_1 = g \circ f_2$, dann gilt $f_1 = f_2$.
- Eine Abbildung $f : M \rightarrow N$ ist surjektiv
 $\iff$ es gibt ein *Rechtsinverses* $e : N \rightarrow M$ mit $f \circ e = \text{id}_N$
 $\iff$ f ist *linkskürzbar*, d. h. wenn $g_1, g_2 : N \rightarrow O$ Abbildungen sind mit $g_1 \circ f = g_2 \circ f$, dann gilt $g_1 = g_2$.
- Eine Abbildung $f : M \rightarrow N$ ist daher bijektiv, wenn sie sowohl ein Rechtsinverses als auch ein Linksinverses hat (und es folgt dann, dass beide die gleiche Abbildung sind).

Für die Existenz der Rechtsinversen surjektiver Abbildungen braucht man das sogenannte Auswahlaxiom (siehe Seite 19).

Abzählbarkeit

Wenn $f : M \rightarrow N$ eine Bijektion ist und M eine endliche Menge, dann hat N genauso viele Elemente wie M . Die Anzahl der Elemente einer Menge heißt auch ihre *Mächtigkeit* (oder *Kardinalität*), und daher definiert man zwei Mengen als *gleichmächtig*, wenn es zwischen ihnen eine Bijektion gibt. Man unterscheidet nun bei unendlichen Mengen solche, die zu der Menge der natürlichen Zahlen gleichmächtig sind, von „größeren“ unendlichen Mengen (wie zum Beispiel der Menge der reellen Zahlen), die *überabzählbar* heißen. Bei der Bedeutung des Begriffes „abzählbar“ gibt es wieder einmal zwei Versionen:

	gleichmächtig zu $\mathbb{N}$	endlich oder gleichmächtig zu $\mathbb{N}$
Version 1:	abzählbar	höchstens abzählbar
Version 2:	abzählbar unendlich	abzählbar

Auf jeder Menge von Mengen ist „gleichmächtig sein“ eine Äquivalenzrelation. Für endliche Mengen wird die Mächtigkeit durch die natürlichen Zahlen gemessen; für beliebige Mengen gibt es dafür die *Kardinalzahlen* als Verallgemeinerung der natürlichen Zahlen.

1.5 Unendliche Konstruktionen

Mengentheoretische Konstruktionen aus unendlich vielen Mengen sind schwieriger zu handhaben, als wenn es nur um endlich viele Mengen geht. Eine *Familie von Mengen* $(M_i)_{i \in I}$ besteht aus einer (Index-)Menge I und einer Zuordnung, die jedem Index $i \in I$ eine Menge M_i zuordnet. Diese Mengen müssen für verschiedene Indizes nicht unbedingt verschieden sein. Die Indexmenge selbst kann endlich oder unendlich sein und auch die leere Menge sein.

Genau genommen ist eine Familie gegeben durch eine Abbildung $f : I \rightarrow M$ für Mengen I und M und $M_i := f(i)$. Es muss also M eine Menge sein, die alle Mengen M_i als Elemente enthält. Gegenüber der oben gegebenen „naiven“ Definition kommt also noch hinzu, dass die Abbildung selbst als Menge existieren muss.

Für solche (möglicherweise unendlichen) Familien gibt es nun ebenfalls Schnitt-, Vereinigungs- und Produktmenge und es gelten die in Tabelle 2 auf Seite 18 aufgeführten Regeln.

- Falls $I \neq \emptyset$, so ist der *Schnitt* $\bigcap_{i \in I} M_i$ der Familie eine Menge, die wie folgt definiert ist:

$$x \in \bigcap_{i \in I} M_i \iff \text{für alle } i \in I \text{ gilt } x \in M_i.$$

Im Folgenden soll immer implizit angenommen sein, dass die Indexmenge eines Schnitts nicht die leere Menge ist.

Nach Definition müsste $\bigcap_{i \in \emptyset} M_i$ alles enthalten, also die „Menge aller Objekte“ sein. Diese gibt es aber nicht, weil es dann auch die „Menge aller Mengen“ geben müsste.

- Die *Vereinigung* $\bigcup_{i \in I} M_i$ der Familie ist eine Menge, die wie folgt definiert ist:

$$x \in \bigcup_{i \in I} M_i \iff \text{es gibt ein } i \in I \text{ mit } x \in M_i.$$

- Das *Produkt* $\prod_{i \in I} M_i$ der Familie ist die Menge der Graphen von Abbildungen $g : I \rightarrow \bigcup_{i \in I} M_i$, welche für alle $i \in I$ die Bedingung $g(i) \in M_i$ erfüllen.

Tabelle 2: REGELN FÜR FAMILIEN VON MENGEN

Aufteilungsregeln	Seien J und K Indexmengen. Dann gilt: $\bigcap_{i \in J \cup K} M_i = \left(\bigcap_{i \in J} M_i \right) \cap \left(\bigcap_{i \in K} M_i \right)$ $\bigcup_{i \in J \cup K} M_i = \left(\bigcup_{i \in J} M_i \right) \cup \left(\bigcup_{i \in K} M_i \right)$ Wenn J und K disjunkt sind, kann man $\prod_{i \in J \cup K} M_i$ mit $\left(\prod_{i \in J} M_i \right) \times \left(\prod_{i \in K} M_i \right)$ identifizieren.
Kommutativgesetze	Wenn $f : I \rightarrow I$ eine Bijektion ist, dann gilt $\bigcap_{i \in I} M_i = \bigcap_{i \in I} M_{f(i)} \quad \text{und} \quad \bigcup_{i \in I} M_i = \bigcup_{i \in I} M_{f(i)}$
Regeln von de Morgan	$M \setminus \left(\bigcap_{i \in I} M_i \right) = \bigcup_{i \in I} (M \setminus M_i)$ $M \setminus \left(\bigcup_{i \in I} M_i \right) = \bigcap_{i \in I} (M \setminus M_i)$
Distributivgesetze	$\left(\bigcap_{i \in I} M_i \right) \cup \left(\bigcap_{j \in J} N_j \right) = \bigcap_{(i,j) \in I \times J} (M_i \cup N_j)$ $\left(\bigcup_{i \in I} M_i \right) \cap \left(\bigcup_{j \in J} N_j \right) = \bigcup_{(i,j) \in I \times J} (M_i \cap N_j)$
Distributivgesetze	$\left(\bigcap_{i \in I} M_i \right) \times N = \bigcap_{i \in I} (M_i \times N) \quad \text{und} \quad N \times \left(\bigcap_{i \in I} M_i \right) = \bigcap_{i \in I} (N \times M_i)$ $\left(\bigcup_{i \in I} M_i \right) \times N = \bigcup_{i \in I} (M_i \times N) \quad \text{und} \quad N \times \left(\bigcup_{i \in I} M_i \right) = \bigcup_{i \in I} (N \times M_i)$
Rechnen mit $\emptyset$	Falls $M_i = \emptyset$ für ein $i \in I \neq \emptyset$, so ist $\prod_{i \in I} M_i = \emptyset$.

Für eine endliche Indexmenge $I = \{1, \dots, n\}$ schreibt man auch $\bigcap_{i=1}^n M_i$ für $\bigcap_{i \in I} M_i$. Falls $I = \{i_1, i_2\}$ zweielementig ist, so stimmt der Schnitt $\bigcap_{i \in I} M_i$ mit dem zuvor definierten Schnitt $M_{i_1} \cap M_{i_2}$ überein. Analog für Vereinigungen.

Wenn I endlich ist, z. B. $I = \{1, \dots, n\}$, sollte $\prod_{i \in I} M_i$ eigentlich nichts anderes sein als das bereits definierte Produkt $M_1 \times \dots \times M_n$. Man würde n -Tupel daher gerne definieren als die Abbildungen von $\{1, \dots, n\}$ in eine geeignete Menge. Dummerweise braucht man in einem mengentheoretischen Aufbau der Mathematik die geordneten Paare, um Abbildungen definieren zu können. Für alle praktischen Belange kann man aber wieder problemlos annehmen, dass $M_1 \times \dots \times M_n = \prod_{i=1}^n M_i$.

Aus dieser Identifikation ergibt es sich auch, dass $\emptyset$ sinnvollerweise ein 0-Tupel (und das einzige) ist.

Für die „leere Vereinigung“ gilt $\bigcup_{i \in \emptyset} M_i = \emptyset$ und für das „leere Produkt“ $\prod_{i \in \emptyset} M_i = \{\emptyset\}$.

Auswahlaxiom, Wohlordnungssatz und das Zornsche Lemma

Das Verhalten unendlicher Mengen lässt sich nicht in allen Fällen aus dem Verhalten endlicher Mengen ableiten. Manchmal braucht man dazu besondere Axiome. Eines davon ist das Auswahlaxiom (mehr zur Stellung des Auswahlaxioms in der Mathematik auf Seite 40).

Auswahlaxiom: Falls $M_i \neq \emptyset$ für alle $i \in I$, so ist $\prod_{i \in I} M_i \neq \emptyset$.

Ist $f \in \prod_{i \in I} M_i$, so gilt $f(i) \in M_i$, d. h. die Abbildung „wählt“ aus jeder der Mengen M_i ein Element aus. Eine äquivalente Formulierung des Auswahlaxioms ist daher: Wenn $(M_i)_{i \in I}$ eine Familie von nicht-leeren Mengen ist, dann gibt es eine Funktion $f : I \rightarrow \bigcup_{i \in I} M_i$ mit $f(i) \in M_i$ für alle $i \in I$.

Die Terminologie ist allerdings etwas unglücklich. Nicht jedesmal, wenn etwas ausgewählt wird, braucht man das Auswahlaxiom, und sehr häufig wird es in der Mathematik benutzt, ohne dass es unmittelbar ersichtlich ist.

Wenn $f : M \rightarrow N$ surjektiv ist, dann ist $(f^{-1}(\{n\}))_{n \in N}$ eine Familie nicht-leerer Mengen mit Indexmenge N . Eine „Auswahlfunktion“ $e \in \prod_{n \in N} f^{-1}(\{n\})$ ist dann ein Rechtsinverses von f , d. h. es gilt $f \circ e = \text{id}_N$.

Mit dem Auswahlaxiom bekommt man auch allgemeinere Distributivgesetze für Familien von Familien von Mengen: $(J_i)_{i \in I}$ ist dabei eine Familie von Indexmengen, und für jedes $i \in I$ ist zudem eine Familie $(M_{i,j})_{j \in J_i}$ gegeben. Dann gilt:

$$\bigcup_{i \in I} \bigcap_{j \in J_i} M_{i,j} = \bigcap_{f \in \prod_{i \in I} J_i} \bigcup_{i \in I} M_{i,f(i)} \quad \text{und} \quad \bigcap_{i \in I} \bigcup_{j \in J_i} M_{i,j} = \bigcup_{f \in \prod_{i \in I} J_i} \bigcap_{i \in I} M_{i,f(i)}$$

Das Auswahlaxiom braucht man, um die Existenz der Funktionen f , über die indiziert wird, sicherzustellen.

Aus dem Auswahlaxiom lassen sich einige für Beweise sehr nützliche Folgerungen ziehen, u. a. den Wohlordnungssatz und das Zornsche Lemma (die beide auf der Basis der anderen Axiome der Mengenlehre äquivalent zum Auswahlaxiom sind). Eine **Wohlordnung** auf einer Menge M ist eine totale Ordnung mit der Eigenschaft, dass jede nicht-leere Teilmenge von M ein kleinstes Element bezüglich dieser Ordnung hat.

Wohlordnungssatz (Zermelo 1904): Auf jeder Menge M gibt es eine Wohlordnung.

Die Nützlichkeit des Wohlordnungssatzes liegt darin, dass man mit Wohlordnungen induktive Beweise führen kann, ähnlich wie mit dem Beweisprinzip der vollständigen Induktion bei den natürlichen Zahlen.

Sei $(M, \leq)$ eine partielle Ordnung.

- Eine *Kette* in $(M, \leq)$ ist eine Teilmenge $K \subseteq M$, die durch $\leq$ total geordnet ist.
- Eine *obere Schranke* einer Teilmenge $A \subseteq M$ ist ein Element von M , das größer als alle Elemente von A ist.
- Ein *maximales Element* von $(M, \leq)$ ist ein Element, zu dem es kein größeres gibt.
(Es kann mehrere maximale Elemente geben. Dagegen ist ein *größtes Element*, also eines, das größer als alle anderen, eindeutig bestimmt, falls es existiert.)

Zornsches Lemma (Kuratowski 1922, Zorn 1935)

Jede nicht-leere, partiell geordnete Menge, in der jede Kette eine obere Schranke hat, enthält mindestens ein maximales Element.

Das Zornsche Lemma ist die Version des Auswahlaxioms, die in der Algebra meist im Umgang mit unendlichen Mengen verwendet wird.

2 Die natürlichen Zahlen und das Induktionsprinzip

Die natürlichen Zahlen sind die Zahlen, mit denen man die Anzahl der Elemente endlicher Mengen beschreiben kann, also die Zahlen $0, 1, 2, 3, \dots$. Meist wird die Menge der natürlichen Zahlen mit $\mathbb{N}$ bezeichnet. In der Mathematik nimmt man nun im allgemeinen an, dass für die natürlichen Zahlen die Eigenschaften gelten, die im Rest dieses Abschnitts beschrieben werden.

Historisch steht das Symbol $\mathbb{N}$ für die Menge der echt positiven Zahlen $1, 2, 3, \dots$ und dann $\mathbb{N}_0$ für die Menge der natürlichen Zahlen einschließlich 0. Auch in einigen Disziplinen wie z. B. den Teilen der Zahlentheorie, welche die natürlichen Zahlen untersuchen, wird die 0 gerne weggelassen, weil sie sich in mancher Hinsicht anders als die anderen natürlichen Zahlen verhält. Ob $\mathbb{N}$ und der Begriff „natürliche Zahl“ also die 0 beinhaltet oder nicht, ist von Autor zu Autor verschieden. Falls die 0 nicht zu $\mathbb{N}$ gezählt wird, wird die Menge $\{0, 1, 2, 3, \dots\}$ gerne die Menge der „nicht negativen ganzen Zahlen“ genannt.

Achtung: Auch bei der Verwendung der Begriffe „positiv“ und „negativ“ herrscht keine Einigkeit: Für manche ist 0 sowohl positiv als auch negativ; sie verwenden dann die Formulierung „echt positiv“ für die Zahlen größer als 0. Für andere ist 0 weder positiv noch negativ; diese sprechen dann von den „nicht negativen ganzen Zahlen“ für die natürlichen Zahlen einschließlich 0.

Auf den natürlichen Zahlen gibt es als grundlegende Struktur die zweistelligen Operationen *Addition* $+$ und *Multiplikation* $\cdot$ und die binäre *Ordnungsrelation* $<$. Addition und Multiplikation sind jeweils kommutativ und assoziativ und zwischen ihnen gilt das Distributivgesetz, d. h. alle natürlichen Zahlen x, y, z erfüllen die Regeln:

$$\begin{aligned} x + y &= y + x & (x + y) + z &= x + (y + z) \\ x \cdot y &= y \cdot x & (x \cdot y) \cdot z &= x \cdot (y \cdot z) \\ (x + y) \cdot z &= (x \cdot z) + (y \cdot z) \end{aligned}$$

Den Multiplikationspunkt $\cdot$ lässt man oft auch weg; wegen der Assoziativgesetze und der Prioritätenregel „Punkt vor Strich“ kann man sich viele Klammern sparen. *Zum Beispiel kann das Distributivgesetz als $(x + y)z = xz + yz$ geschrieben werden.*

Die binäre Relation $<$ ist eine strikte totale Ordnung, die mit der Addition und der Multiplikation in der folgenden Weise verträglich ist. Für alle natürlichen Zahlen x, y, z gilt:

$$\begin{aligned} \text{Wenn } x < y, \text{ dann auch } x + z < y + z. \\ \text{Wenn } x < y \text{ und } z \neq 0, \text{ dann auch } xz < yz. \end{aligned}$$

Die Ordnung $<$ auf den natürlichen Zahlen hat ein Minimum 0 und kein Maximum (d. h. es gibt keine größte natürliche Zahl). Außerdem ist es eine *diskrete Ordnung*, d. h. zu jeder natürlichen Zahl gibt es einen *Nachfolger* und zu jeder natürlichen Zahl außer 0 einen *Vorgänger*. Ein Nachfolger von x ist eine kleinste Zahl, die größer als x ist, und ein Vorgänger eine größte Zahl, die kleiner als x ist. Da die Ordnung total ist, sind Nachfolger und Vorgänger eindeutig bestimmt, und der Nachfolger von x ist natürlich $x + 1$, der Vorgänger von $x \neq 0$ ist $x - 1$.

Außerdem ist die Ordnungsrelation auf den natürlichen Zahlen eine sogenannte *Wohlordnung*. Das bedeutet, dass jede nicht-leere Teilmenge der natürlichen Zahlen ein kleinstes Element hat. Auf dieser Eigenschaft beruht das *Induktionsprinzip*:

Jede Teilmenge der natürlichen Zahlen, welche die Zahl 0 enthält und mit jeder Zahl auch ihren Nachfolger, ist gleich der Menge aller natürlichen Zahlen.

Die (naiv gebildete) Teilmenge der natürlichen Zahlen, die aus 0 besteht, dem Nachfolger von 0, dem Nachfolger des Nachfolgers von 0 usw. ist nach diesem Prinzip bereits die Menge aller natürlichen Zahlen, d. h. jede natürliche Zahl ergibt sich in endlich vielen Schritten aus der 0 durch Anwenden der Nachfolgeroperation. In diesem Sinne besagt das Induktionsprinzip also, dass alle natürlichen Zahlen endlich sind und es keine unendlich großen natürlichen Zahlen gibt.

Tatsächlich gibt es hier eine Schwierigkeit, die in der mathematischen Logik eine große Rolle spielt. Die naheliegenden Definitionen von „endlich“ und „natürlicher Zahl“ bedingen sich gegenseitig: Eine Menge ist endlich, wenn die Anzahl ihrer Elemente eine natürliche Zahl ist, und die natürlichen Zahlen sind diejenigen, die man in endlich vielen Schritten aus der Null durch die Nachfolgeroperation gewinnt. Das Induktionsprinzip ist ein axiomatischer Ausweg aus diesem Teufelskreis: Es beschreibt die geforderte Eigenschaft, ohne das Wort „endlich“ zu benutzen.

2.1 Beweis per vollständiger Induktion

Das Induktionsprinzip ermöglicht das Beweisverfahren durch vollständige Induktion, das es in mehreren Varianten gibt. Man möchte zeigen, dass alle natürlichen Zahlen eine Eigenschaft E haben.

Vollständige Induktion, Variante 1:

Induktionsanfang: Man zeigt, dass die Zahl 0 die Eigenschaft E hat.

Induktionsschritt: Man zeigt: Wenn eine beliebige natürliche Zahl x die Eigenschaft E hat, dann hat auch ihr Nachfolger $x + 1$ die Eigenschaft E .

Vollständige Induktion, Variante 2:

Induktionsanfang: Man zeigt, dass (für eine fest vorgegebene natürliche Zahl k) die Zahlen $0, \dots, k$ die Eigenschaft E haben.

Induktionsschritt: Man zeigt: Wenn für eine beliebige natürliche Zahl x die Zahlen $x, x + 1, \dots, x + k$ die Eigenschaft E haben, dann hat auch die Zahl $x + k + 1$ die Eigenschaft E .

Vollständige Induktion, Variante 3:

Induktionsanfang: Man zeigt, dass die Zahl 0 die Eigenschaft E hat.

Induktionsschritt: Man zeigt: Wenn für eine beliebige natürliche Zahl x alle natürlichen Zahlen $0, \dots, x$ die Eigenschaft E haben, dann hat auch die Zahl $x + 1$ die Eigenschaft E .

Das Induktionsprinzip liefert dann jeweils das gewünschte Ergebnis, nämlich dass alle natürlichen Zahlen die Eigenschaft E haben.

Den Beweis per vollständiger Induktion gibt es auch in einer Version als Widerspruchsbeweis:

Methode des kleinsten Verbrechers:

Man zeigt, dass die Zahl 0 die Eigenschaft E hat.

Man nimmt an, dass n die kleinste natürliche Zahl ist, die nicht die Eigenschaft E hat, und zeigt, dass es eine kleinere Zahl geben muss, welche die Eigenschaft E hat.

Eine gängige Sprechweise lautet „Induktion nach n “, wenn man deutlich machen will, dass in einer Aussagen mit eventuell mehreren natürlichen Zahlen die Variable n diejenige ist, auf die man das Induktionsprinzip anwendet. Mit diesen Beweisprinzipien lassen sich auch Eigenschaften für alle natürlichen Zahlen ab einer festen Zahl n_0 zeigen, wenn man den Induktionsanfang jeweils bei n_0 statt bei 0 ansetzt.

Oft wird das Induktionsprinzip auf eine indirekte Weise angewandt. Gegeben ist dann eine Menge M und eine Abbildung $f : M \rightarrow \mathbb{N}$. Man zeigt dann „per Induktion nach $f(x)$ “, z. B. in [Anwendung von Variante 3](#), dass alle $x \in M$ eine Eigenschaft E haben, indem man beweist:

Wenn für eine beliebige Zahl n alle $x \in M$ mit $f(x) < n$ die Eigenschaft E haben, dann haben auch alle $x \in M$ mit $f(x) = n$ die Eigenschaft E .

(In dieser Version sind übrigens Induktionsanfang und -schritt zusammengefasst: für $n = 0$ ergibt sich der Induktionsanfang, für $n > 0$ der Induktionsschritt.)

Per Induktion kann man auch Definitionen einführen. Solche Definitionen werden auch *rekursive Definitionen* genannt.

Zum Beispiel wird die *Fibonacci-Folge* $(F_n)_{n \in \mathbb{N}}$ üblicherweise rekursiv definiert durch die Anfangswerte $F_0 := 0$ und $F_1 := 1$ (die dem Induktionsanfang entsprechen) und die Rekursionsformel $F_n := F_{n-2} + F_{n-1}$ für alle $n \geq 2$ (die dem Induktionsschritt entspricht). Hier ist die Definition analog zum Induktionsverfahren in der Variante 2 mit $k = 1$.

2.2 Summen und Produkte

Ähnlich wie es für Schnitt, Vereinigung und Mengenprodukt die iterierten zweistelligen Operationen mit „kleinen Symbolen“ $\cap, \cup, \times$ und die Operationen auf Familien mit „großen Symbolen“ $\bigcap, \bigcup, \prod$ gibt, kann man auch Summen und Produkte von „Familien von Zahlen“ mit „großen Symbolen“ schreiben.

Dazu gibt es das *Summenzeichen* $\sum$ und das *Produktzeichen* $\prod$.

Das Summenzeichen ist ein großes Sigma (griechischer Buchstabe für „S“ wie Summe) und das Produktzeichen ein großes Pi (griechischer Buchstabe für „P“ wie Produkt).

Die Doppelverwendung des Produktzeichens für Mengen- und für Zahlenprodukte führt normalerweise nicht zu Verwirrung.

Diese Symbole werden in folgenden Varianten benutzt, hier z. B. mit der Indexmenge $\{1, \dots, n\}$:

$$\begin{aligned} \sum_{i=1}^n a_i &= \sum_{i=1}^n a_i = \sum_{i \in \{1, \dots, n\}} a_i = \sum \{a_i \mid i \in \{1, \dots, n\}\} = a_1 + a_2 + \dots + a_n \\ \prod_{i=1}^n a_i &= \prod_{i=1}^n a_i = \prod_{i \in \{1, \dots, n\}} a_i = \prod \{a_i \mid i \in \{1, \dots, n\}\} = a_1 \cdot a_2 \cdot \dots \cdot a_n \end{aligned}$$

Analog für andere Indexmengen, z. B. $\sum_{i=5}^8 a_i = a_5 + a_6 + a_7 + a_8$. Die Indexmengen können auch aus andere diskreten Ordnungen stammen, z. B. den ganzen Zahlen. So ist $\sum_{i=-3}^2 a_i = a_{-3} + a_{-2} + a_{-1} + a_0 + a_1 + a_2$.

Im Prinzip gelten diese Schreibweisen auch für unendliche Indexmengen. Allerdings muss erst einmal definiert werden, was eine unendliche Summe bzw. ein unendliches Produkt bedeuten soll. Dies geschieht in Analysis I. *Sofern die Ausdrücke überhaupt definiert sind*, kann man sie analog zum endlichen Fall auf die folgenden Weisen schreiben, hier am Beispiel der Indexmenge $\mathbb{N}$:

$$\begin{aligned} \sum_{i=0}^{\infty} a_i &= \sum_{i=0}^{\infty} a_i = \sum_{i \in \mathbb{N}} a_i = \sum \{a_i \mid i \in \mathbb{N}\} = a_0 + a_1 + \dots + a_n + \dots \\ \prod_{i=0}^{\infty} a_i &= \prod_{i=0}^{\infty} a_i = \prod_{i \in \mathbb{N}} a_i = \prod \{a_i \mid i \in \mathbb{N}\} = a_0 \cdot a_1 \cdot \dots \cdot a_n \cdot \dots \end{aligned}$$

Entsprechende Verwendung finden die Symbole für andere unendliche Indexmenge, sofern eben die unendliche Summe bzw. das unendliche Produkt über diese Indexmenge überhaupt definiert ist.

Auch hier gibt es die Sonderfälle „leerer Summen“ und „leerer Produkte“. Es ist $\sum_{i \in \emptyset} a_i = 0$ und $\prod_{i \in \emptyset} a_i = 1$. Ebenso gilt dies für die Schreibweisen $a_1 + \dots + a_n$ bzw. $a_1 \cdot \dots \cdot a_n$ im Fall $n = 0$: es ist dann $a_1 + \dots + a_0 = 0$ und $a_1 \cdot \dots \cdot a_0 = 1$. Das intuitive Verständnis hierfür ist, dass man eine leere Summe zu einer bestehenden Summe anfügen können sollte, ohne dass sich etwas ändert. Dazu muss die leere Summe das neutrale Element der Addition sein, also die Null. Analog ist das leere Produkt das neutrale Element der Multiplikation, also die Eins.

Die Summen- und Produktschreibweisen mit $\sum$ und $\prod$ werden auch in anderen Situationen angewandt, in denen man eine kommutative und assoziative Addition bzw. Multiplikation hat.

3 Logik (oder: „Wie redet man in der Mathematik?“)



WARNUNG III



Dieser Abschnitt kann wegen seiner Kürze das Thema nur anreißen und bleibt daher an vielen Stellen im Vagen. Für genaue Definitionen und ein besseres Verständnis sollte man eine Vorlesung in Mathematischer Logik hören.

Mathematische Sachverhalte und Theoreme werden als Aussagesätze formuliert, und auch das mathematische Argumentieren und Beweisen erfolgt über Aussagesätze bzw. Aussagen.² Für den Umgang mit diesen Aussagen legt man gewisse Regeln zugrunde, „eine Logik“. Hauptsächlich geht es dabei um zwei Punkte:

- Begriffe wie „Logische Folgerung“ und „Logische Äquivalenz“ müssen definiert werden.
- Die Bedeutung eines zusammengesetzten Aussagesatzes aus seinen Teilen muss erklärt werden (zum Beispiel um verstehen zu können, wie die Verneinung einer sprachlich komplexen Aussage aussieht).

Es gibt einen Teil der Logik, der das Verhältnis von Aussagesätzen zueinander klärt, die sogenannte *Aussagenlogik*, und einen weitergehenden Teil, der die Struktur von Aussagen näher analysiert, insbesondere in Hinblick auf *Quantifizierungen*, dies ist dann die *Prädikatenlogik*. Üblicherweise wird in der Mathematik, zumindest im mathematischen Argumentieren innerhalb von Beweisen, die im Folgenden erläuterte *klassische zweiwertige Logik* benutzt. In einigen Punkten weicht sie vom Alltagsgebrauch ab und einige ihrer Regeln entsprechen nicht der Intuition. An beides muss man sich gewöhnen.

Welche Logik man benutzt, ist eine Frage der Konvention und nicht eine Frage einer höheren Wahrheit. Die klassische Aussagenlogik ist für die Mathematik üblich und sinnvoll, im Alltag aber nicht üblich und daher auch nicht immer sinnvoll. Wenn man zum Beispiel einem Kind sagt: „Wenn du deine Hausaufgaben nicht machst, darfst du nicht fernsehen“ ist bei einem streng logischen Verständnis des Satz nichts über den Fall gesagt, dass das Kind seine Hausaufgaben macht. Das Kind wird den Satz aber auch als Implikation in die andere Richtung verstehen: „Wenn du deine Hausaufgaben machst, darfst du fernsehen“. Ein Fernsehverbot trotz Hausaufgaben ist dann zwar kein logischer Fehler, führt aber trotzdem zu Geschrei und Familienkrach.

Es gibt eine kleine Minderheit von Mathematikern, die strengere Anforderungen an mathematische Beweise stellen und mit schwächeren Logiken arbeiten, etwa dem Intuitionismus (die schwächer in dem Sinne ist, dass weniger Regeln gelten, z. B. gilt das Prinzip des ausgeschlossenen Dritten nicht mehr).

Der klassischen Aussagenlogik liegen die beiden folgenden Prinzipien zugrunde:

- *Zweiwertigkeit (mit den Prinzipien des ausgeschlossenen Widerspruchs und des ausgeschlossenen Dritten):*

Um das Verhalten von Aussagen zu verstehen, ist es ausreichend, stets nur Situationen zu betrachten, in denen jede Aussage *entweder* wahr *oder* falsch ist (auch wenn man vielleicht nicht in der Lage ist festzustellen, welcher von beiden Fällen zutrifft). Insbesondere kann es also nicht sein, dass eine Aussage gleichzeitig wahr und falsch ist (*ausgeschlossener Widerspruch*), und es kann nicht sein, dass eine Aussage weder wahr noch falsch ist (*ausgeschlossenes Drittes*). Die beiden Möglichkeiten „wahr“ und „falsch“ werden die möglichen *Wahrheitswerte* der Aussage genannt.

- *Kompositionalität und Kontextfreiheit:*

Der Wahrheitswert einer aus Teilaussagen zusammengesetzten Aussage hängt nur von den Wahrheitswerten der Teilaussagen und der Art der Zusammensetzung ab.

²Für das, was folgt, ist es nicht nötig, zwischen Aussagesätzen und den von ihnen gemachten Aussagen zu unterscheiden.

Es folgen nun die wichtigsten in der Mathematik benutzten Arten der Zusammensetzung von Aussagen aus Teilaussagen. Um sie zu verstehen, muss man wegen des Kompositionalitätsprinzips also nur wissen, wie sich der Wahrheitswert der Gesamtaussage aus den Wahrheitswerten der Teilaussagen ergibt.

Einstellige Zusammensetzungen:

- Die **Negation** „Es ist nicht der Fall, dass A “ oder kurz: „*nicht* A “ einer Aussage A , mit folgendem Wahrheitswertverhalten:

A	wahr	falsch
nicht A	falsch	wahr

Zweistellige Zusammensetzungen:

- Die **Konjunktion** „ A und B “ zweier Aussagen A und B , mit folgendem Wahrheitswertverhalten:

A	wahr	wahr	falsch	falsch
B	wahr	falsch	wahr	falsch
A und B	wahr	falsch	falsch	falsch

- Die **Disjunktion** „ A oder B “ zweier Aussagen A und B , mit folgendem Wahrheitswertverhalten:

A	wahr	wahr	falsch	falsch
B	wahr	falsch	wahr	falsch
A oder B	wahr	wahr	wahr	falsch

Insbesondere ist mit „oder“ in der Mathematik stets das einschließende Oder gemeint, im Gegensatz zu dem ausschließenden „*entweder ... oder ...*“.

- Die **Implikation** „wenn A , dann B “ zweier Aussagen A und B , mit folgendem Wahrheitswertverhalten:

A	wahr	wahr	falsch	falsch
B	wahr	falsch	wahr	falsch
wenn A , dann B	wahr	falsch	wahr	wahr

- Die **Äquivalenz** „ A genau dann, wenn B “ zweier Aussagen A und B , mit folgendem Wahrheitswertverhalten:

A	wahr	wahr	falsch	falsch
B	wahr	falsch	wahr	falsch
A genau dann, wenn B	wahr	falsch	falsch	wahr

Bei verschachtelten Zusammensetzungen muss man auf die Reihenfolge achten, die umgangssprachlich manchmal schwer auszudrücken ist. Zum Beispiel muss man „es ist nicht der Fall, dass A , und B “ und „es ist nicht der Fall, dass A und B “ unterscheiden. Ich schreibe der Klarheit halber Klammern, also im Beispiel „(nicht A) und B “ bzw. „nicht (A und B)“.

3.1 Logische Folgerung und logische Äquivalenz

Um über das logische Verhalten von Aussagen zu sprechen, benutzt man die folgenden Begriffe, deren genaue Definition leider nur im Rahmen der Entwicklung von formalen Sprachen gegeben werden kann. Ich spreche bei diesen Definitionen ziemlich vage von „allen möglichen Umständen“. Wenn es sich um rein aussagenlogische Verhältnisse handelt, bedeutet dies, dass man

sich jede mögliche Verteilung der Wahrheitswerte für die Teilaussagen anschaut. Später, in der Prädikatenlogik, wird man noch alle möglichen Auswertungen der Quantoren hinzunehmen.

- Eine Aussage B *folgt (logisch) aus* einer Aussage A (auch: „ A *impliziert* B “), wenn folgendes gilt:

Unter allen möglichen Umständen, unter denen die Aussage A wahr wird, wird auch die Aussage B wahr. Oder kurz: Jedesmal, wenn A wahr wird, wird auch B wahr.

Beispiel: Stehen C und D für Aussagen, so folgt „ C *impliziert* D “ aus „*nicht* C “. Hier ist über den Aufbau von C und D nichts bekannt, die „möglichen Umstände“ sind daher die möglichen Wahrheitswertverteilungen für C und D , die man am besten in einer Wahrheitstafel auswertet:

C	D	$\text{nicht } C$	$C \text{ impliziert } D$	
wahr	wahr	falsch	wahr	
wahr	falsch	falsch	falsch	
falsch	wahr	wahr	wahr	←
falsch	falsch	wahr	wahr	←

Es gibt hier vier mögliche Wahrheitswertverteilungen für C und D , von denen zwei „*nicht* C “ wahr werden lassen (in der Tabelle mit einem Pfeil markiert). Unter diesen beiden Verteilungen wird auch „ C *impliziert* D “ wahr.

- Zwei Aussagen A und B sind *(logisch) äquivalent*, wenn sie unter allen möglichen Umständen den gleichen Wahrheitswert annehmen.

„ A “ und „*nicht nicht* A “ sind äquivalent.

- Eine Aussage ist eine *Tautologie*, wenn sie unter allen möglichen Umständen den Wahrheitswert „wahr“ annimmt. Oder kurz: die Aussage ist immer wahr.

Steht A für eine Aussage, so ist „ A *oder nicht* A “ eine Tautologie. A kann wahr oder falsch sein, aber in beiden Fällen wird die Gesamtaussage wahr.

Die Definition bleibt auch im rein aussagenlogischen Fall noch vage, weil nicht genau definiert ist, was überhaupt eine Teilaussage ist und was erlaubte Zusammensetzungen. In der mathematischen Logik wird dies im Rahmen von formalen Sprachen völlig präzise definiert. Vorläufig reicht die Vorstellung, dass alle Zusammensetzungen betrachtet werden, die „wahrheitswertfunktional“ sind, d. h. zu deren Verständnis man nur wissen muss, wie sich der Wahrheitswert der Gesamtaussage aus den Teilaussagen errechnet. Man kann übrigens mit den oben angegebenen Zusammensetzungen auskommen.

Zwei Dinge sind besonders zu beachten:

- ☞ Die Definition von „logischer Folgerung“ hat selbst die Form einer Implikation (ersichtlich in der Kurzversion) und ist selbst bereits im Sinne der oben gegebenen Wahrheitstafel zu verstehen. Das „*Wenn ... dann ...*“ der klassischen Aussagenlogik sagt also nur etwas über das Auftreten möglicher Wahrheitswertverteilungen aus und nichts über einen kausalen oder inneren Zusammenhang der beiden Aussagen.

In der Sprache der Mathematik sind Sätze wie der folgende daher richtig: „*Wenn der Satz von Pythagoras gilt, dann ist 2 eine gerade Zahl*“, weil beide Teilaussagen wahr sind. Im normalen Sprachgebrauch würde man solch eine Aussage eher als unsinnig oder gar falsch ansehen, (a) weil es keinen erkennbaren inneren Zusammenhang zwischen den beiden Teilaussagen gibt und (b) weil der Satz von Pythagoras ja gilt, also keine echte Bedingung darstellt, die auch falsch werden könnte.

- ☞ In der Regel betrachtet man nicht wirklich „alle möglichen Umstände“, sondern nur diejenigen, die mit einem grundlegenden Verständnis der Mathematik verträglich sind. (In der Sprache der Logik: man arbeitet „relativ zu einer Hintergrundtheorie“). Im Beispiel oben zieht man etwa für die Teilaussage „*2 ist eine gerade Zahl*“ gar nicht in Betracht, dass sie

falsch sein könnte, da sie im üblichen mathematischen Verständnis von „2“ und „gerade Zahl“ stimmt, d. h. beim (expliziten oder impliziten) Auswerten von Wahrheitstafeln lässt man die Fälle, dass dieser Aussage der Wahrheitswert *falsch* zukommt, einfach weg.

Nicht immer wird explizit gemacht, auf welche „Hintergrundtheorie“ man sich gerade bezieht, sondern dies muss aus dem Kontext erschlossen werden.

Wenn etwa klar ist, dass man in den reellen Zahlen arbeitet, wird man für die Aussage „Das Polynom $X^2 + 1$ hat eine Nullstelle“ nur den Wahrheitswert *falsch* betrachten. Arbeitet man dagegen in den komplexen Zahlen, wird man für diese Aussage nur den Wahrheitswert *wahr* betrachten. Ist nicht klar, auf welchen Zahlbereich sie sich bezieht oder sind es wechselnde Zahlbereiche, muss man beide Wahrheitswerte betrachten.

Wenn B aus A folgt, sagt man auch, dass A eine *hinreichende Bedingung* für B ist und B eine *notwendige Bedingung* für A . Zwei Aussagen A und B sind also äquivalent, wenn A sowohl notwendige als auch hinreichende Bedingung für B ist.

Bei umgangssprachlichen Ausdrucksweisen muss man manchmal darauf achten, ob eine hinreichende oder eine notwendige Bedingung gemeint ist. Im mathematischen Sprachgebrauch steht die Formulierung „ A gilt nur dann, wenn B “ für eine notwendige Bedingung, also für „aus A folgt B “. Für die Äquivalenz zweier Aussagen A und B sind die folgenden Formulierungen üblich: „ A genau dann, wenn B “ oder „ A dann und nur dann, wenn B “ oder „wenn A und nur wenn A , dann B “ oder ähnliches.

Alle möglichen zusammengesetzten Wahrheitswertverläufe lassen sich übrigens bis auf Äquivalenz aus den oben angegebenen Zusammensetzungen konstruieren, z. B. ist die Aussage „weder A noch B noch C “ äquivalent zu „(nicht A) und (nicht B) und (nicht C)“.

3.2 Symbolik

In der mathematischen Logik werden Aussagen wiederum als mathematische Objekte betrachtet und formalisiert. Der Begriff z. B. der logischen Äquivalenz wird dann als zweistellige Relation auf einer Menge von Sätzen in einer formalen Sprache (sogenannte „logische Formeln“) definiert. Um diese Formeln zu definieren, braucht man Symbole für die Art der Zusammensetzung von Aussagen, sogenannte *Junktoren*. Für diese Symbole gibt es keinen Standard, aber die folgende Version dürfte die in der Mathematik verbreitetste sein:

$\neg A$	für die Negation eines Aussagesatzes A
$(A \wedge B)$	für die Konjunktion der Aussagesätze A und B
$(A \vee B)$	für die Disjunktion der Aussagesätze A und B
$(A \rightarrow B)$	für die Implikation zwischen den Aussagesätzen A und B
$(A \leftrightarrow B)$	für die Äquivalenz der Aussagesätze A und B

Diese Symbole dienen eigentlich nur dazu, die Struktur eines Aussagesatzes darzustellen; sie machen keine Aussage darüber, ob der Satz gilt oder nicht. Allerdings sind sie in den mathematischen Alltag hineingedrungen und werden entgegen der ursprünglichen Intention bisweilen auch als Abkürzungen in mathematischen Aufschrieben benutzt. Anfänger neigen manchmal dazu, sie im Übermaß zu benutzen. Man sollte sie nur dann benutzen, wenn durch sie die Struktur einer Aussage klarer wird, als wenn man sie mit Worten ausdrückt.

Darüberhinaus gibt es die beiden sehr häufig verwendeten abkürzenden Symbole

$$\implies \quad \text{und} \quad \iff,$$

die auf zwei verschiedene Arten verwendet werden.

Die erste (und eigentliche) Verwendungsweise: Hierbei steht die Schreibweise „ $A \implies B$ “ für eine Aussage über die beiden Aussagesätze A und B , nämlich für „die Aussage B folgt aus der Aussage A “. Die Schreibweise „ $A \iff B$ “ steht entsprechend für „die Aussagen A und B sind äquivalent“.

Diese beiden Symbole machen also Aussagen über die Gültigkeit von Aussagesätzen (allerdings in dieser Verwendungsweise nicht darüber, ob A und B gelten oder nicht!). Im Gegensatz dazu wird $(A \rightarrow B)$ benutzt um auszudrücken, dass *ein* Aussagesatz vorliegt, der die Struktur einer Implikation hat (und entsprechend $(A \leftrightarrow B)$ für einen Satz, der die Struktur einer Äquivalenz hat).

Die zweite (uneigentliche) Verwendungsweise: Der Pfeil $\Rightarrow$ wird sehr häufig als Abkürzung für eine Schlussweise benutzt, d. h. man schreibt:

$$\begin{array}{c} A \\ \Rightarrow B \end{array} \quad \text{oder auch} \quad A \Rightarrow B$$

um zu sagen: „ A gilt (weil es ein Axiom oder eine Annahme ist oder gerade bewiesen wurde). Man weiß oder sieht leicht ein, dass B aus A folgt. Also gilt auch B .“

Analog für den Doppelpfeil $\Leftrightarrow$.

Tabelle 3: AUSSAGENLOGISCHE GESETZE

A, B, C sind beliebige Aussagen:

<i>Doppelnegationsregel</i>	$A \Leftrightarrow \neg\neg A$	
<i>Regeln von de Morgan</i>	$\neg(A \wedge B) \Leftrightarrow (\neg A \vee \neg B)$ $\neg(A \vee B) \Leftrightarrow (\neg A \wedge \neg B)$	
<i>Kommutativgesetz für $\wedge$</i>	$(A \wedge B) \Leftrightarrow (B \wedge A)$	
<i>Kommutativgesetz für $\vee$</i>	$(A \vee B) \Leftrightarrow (B \vee A)$	
<i>Assoziativgesetz für $\wedge$</i>	$((A \wedge B) \wedge C) \Leftrightarrow (A \wedge (B \wedge C))$	
<i>Assoziativgesetz für $\vee$</i>	$((A \vee B) \vee C) \Leftrightarrow (A \vee (B \vee C))$	
<i>Distributivgesetze für $\wedge$ und $\vee$</i>	$((A \wedge B) \vee C) \Leftrightarrow ((A \vee C) \wedge (B \vee C))$ $((A \vee B) \wedge C) \Leftrightarrow ((A \wedge C) \vee (B \wedge C))$	
<i>Kontrapositionsregeln</i>	$(A \rightarrow B) \Leftrightarrow (\neg B \rightarrow \neg A)$ $(A \leftrightarrow B) \Leftrightarrow (\neg B \leftrightarrow \neg A)$	
<i>(Anti-)Distributivgesetze für $\rightarrow$ und $\wedge$ bzw $\vee$</i>	$(A \rightarrow (B \wedge C)) \Leftrightarrow ((A \rightarrow B) \wedge (A \rightarrow C))$ $(A \rightarrow (B \vee C)) \Leftrightarrow ((A \rightarrow B) \vee (A \rightarrow C))$ $((A \wedge B) \rightarrow C) \Leftrightarrow ((A \rightarrow C) \vee (B \rightarrow C))$ $((A \vee B) \rightarrow C) \Leftrightarrow ((A \rightarrow C) \wedge (B \rightarrow C))$	$[\text{sic!}]^3$ $[\text{sic!}]$
<i>Definition von $\rightarrow$</i>	$(A \rightarrow B) \Leftrightarrow (\neg A \vee B)$	
<i>Regeln für $\rightarrow$</i>	$(A \rightarrow (B \rightarrow C)) \Leftrightarrow (B \rightarrow (A \rightarrow C))$ $\Leftrightarrow ((A \wedge B) \rightarrow C)$ $\Leftrightarrow ((A \rightarrow B) \rightarrow (A \rightarrow C))$	
<i>Definition von $\leftrightarrow$</i>	$(A \leftrightarrow B) \Leftrightarrow ((A \rightarrow B) \wedge (B \rightarrow A))$	
<i>Substitutionsregel</i>	Man kann in einer logischen Äquivalenz (einer logischen Folgerung, einer Tautologie) eine Teilaussage durch eine äquivalente Aussage ersetzen, und es bleibt eine Äquivalenz (bzw. Folgerung, bzw. Tautologie).	

„Sic!“ bedeutet: es ist kein Tippfehler, obwohl man das glauben könnte.

In der zweiten Verwendungsweise kann man auch problemlos Pfeile hintereinanderschreiben, also etwa

$$A \implies B \implies C,$$

und meint damit: A gilt, also auch B , also auch C . In der ersten Verwendungsweise wäre die Bedeutung nicht klar, da $(A \implies B) \implies C$ und $A \implies (B \implies C)$ verschiedene Aussagen sind.

In der Regel wird aus dem Kontext klar, welche Verwendungsweise vorliegt; es kommt nur äußerst selten zu Verwirrung. Es empfiehlt sich aber, den Gebrauch der Folgerungspfeile $\implies$ zugunsten von Wörtern wie *also*, *folglich*, *mithin* ... einzuschränken.

Ein Beweis wird jedenfalls nicht dadurch richtig, dass viele Folgerungs- oder Äquivalenzpfeile dabei stehen. Die Folgerungspfeile dienen auf keinen Fall dazu, den jeweils nächsten Beweisschritt anzuzeigen.

Das Hintereinanderschreiben von Aussagen bedeutet üblicherweise deren Konjunktion, z. B. steht

$$A \iff B, C$$

dafür, dass die Aussage A genau dann gilt, wenn die Aussage B und die Aussage C gilt.

3.3 Wichtige Regeln

Ob ein (aussagenlogischer) Schluss korrekt ist oder nicht oder ob eine Äquivalenz gilt oder nicht, lässt sich durch das Ausrechnen der Wahrheitswertverläufe nachprüfen. Nicht alle gültigen Gesetze sind immer intuitiv einsichtig (z. B. $\neg A \iff (A \rightarrow \neg A)$).

Zur Verdeutlichung: Eine Regel wie $\neg A \iff (A \rightarrow \neg A)$ bedeutet: Wenn A eine Aussage ist, B eine Aussage von der Form „nicht A “ und C eine Aussage von der Form „wenn A , dann nicht A “, dann sind B und C logisch äquivalent.

In Tabelle 3 auf Seite 28 sind einige grundlegende Regeln in dieser Formelschreibweise angegeben.

Die Parallelität der Regeln für Konjunktion und Disjunktion einerseits mit den Regeln für Schnitt und Vereinigung andererseits erklärt sich daraus, dass bei der Definition von Schnitt bzw. Vereinigung die Konjunktion bzw. Disjunktion benutzt wurde. Diese Parallelität hat auch eine mathematische Formulierung: Die Potenzmenge $\mathfrak{P}(M)$ einer Menge M bildet mit den beiden zweistelligen Operationen $\cap$ und $\cup$ eine sogenannte *Boolesche Algebra*. Die Aussagen (bis auf logische Äquivalenz) bilden mit den Operationen „Konjunktion“ und „Disjunktion“ in der klassischen Aussagenlogik ebenfalls eine Boolesche Algebra.

Hier noch eine Mahnung zur Vorsicht: Stehen A und B für Aussagen, so ist nach der klassischen Aussagenlogik „ $((A \rightarrow B) \vee (A \rightarrow \neg B))$ “ eine Tautologie. Dies bedeutet aber nicht, dass man deshalb aus der Aussage A die Aussage B oder die Aussage „nicht B “ folgern könnte (sondern lediglich die Tautologie „ B oder nicht B “). Man könnte dies so formulieren, dass die logische Verknüpfung „oder“ nicht mit den Anführungszeichen vertauscht, oder dass das Zeichen $\vee$ in der Formel nicht dem Wort „oder“ im dritten Satz dieser Bemerkung entspricht. Manche halten dies für einen Konstruktionsfehler der klassischen Aussagenlogik und lehnen sie daher ab.

3.4 Beweise und logische Schlüsse

Hier folgt nun noch ein wenig Hintergrund zur zweiten Verwendungsweise der Folgerungs- bzw. Äquivalenzpfeile:

Eine mathematische Argumentation (ein Beweis) lässt sich in kleine Schritte aufteilen, in denen *logische Schlüsse* gezogen werden. Solch ein Schluss besteht zum einen aus Aussagen (den *Prämissen* des Schlusses), deren Gültigkeit man bereits einsieht (weil man sie bereits bewiesen hat, oder weil es Axiome sind, oder Annahmen, unter denen der Beweisschritt stattfindet), und zum andern aus einer Aussage (der *Konklusion*), die aus den Prämissen gefolgert wird. Der Schluss ist korrekt, wenn die Konklusion im Sinne der obigen Definition logisch aus der Konjunktion der Prämissen folgt, wenn also die Gültigkeit der Prämissen auf die Gültigkeit der Konklusion

schließen lässt, d. h. wenn es nicht möglich ist, dass alle Prämissen wahr sind und die Konklusion falsch. (Dies ist auch die Begründung dafür, dass die Wahrheitstafel für die Implikation, so wie sie oben gegeben wurde, sinnvoll ist.)

Beispiel eines korrekten Schlusses („Fallunterscheidung“ genannt):

$$\begin{array}{ll} \text{Prämisse} & (A \rightarrow B) \\ \text{Prämisse} & (\neg A \rightarrow B) \\ \hline \text{Konklusion} & B \end{array}$$

Der elementarste aller korrekten Schlüsse ist der *Modus Ponens*:

$$\begin{array}{ll} \text{Prämisse} & A \\ \text{Prämisse} & (A \rightarrow B) \\ \hline \text{Konklusion} & B \end{array}$$

Der Folgerungspfeil $\Rightarrow$ steht also in seiner zweiten Verwendungsweise als Abkürzung für einen Modus Ponens.

Mit Wahrheitstafeln kann man leicht die Korrektheit eines Schlusses nachweisen. Die meisten der üblicherweise verwendeten elementaren Schlüsse sind auch sehr einsichtig. Nur zwei Schlussweisen, die am Anfang öfters etwas Mühe bereiten, sollen näher beschrieben werden:

Kontraposition etc.

Angenommen man weiß, dass B aus A folgt (also dass die Implikation $(A \rightarrow B)$ immer wahr ist). Mit der Kontrapositionsregel hat man dann auch, dass „nicht A “ aus „nicht B “ folgt (also dass auch $(\neg B \rightarrow \neg A)$ immer wahr ist).

- Wenn man zusätzlich weiß, dass A gilt, kann man auf die Gültigkeit von B schließen (das ist der *Modus Ponens*).
- Wenn man zusätzlich weiß, dass „nicht B “ gilt, kann man auf die Gültigkeit von „nicht A “ schließen (dieser Schluss heißt *Modus Tollens*; es ist der auf die Kontraposition angewandte Modus Ponens).
- Dagegen: Wenn man zusätzlich weiß, dass B gilt, kann man *nicht* auf die Gültigkeit von A schließen. Ebenso: Wenn man zusätzlich weiß, dass „nicht A “ gilt, kann man *nicht* auf die Gültigkeit von „nicht B “ schließen.

Hier folgen nun (in Symbolschreibweise) die acht korrekten Schlussweisen, die sich als Varianten des Modus Ponens bzw. Tollesn ergeben, wenn auch noch einfache Negationen ins Spiel kommen:

$$\begin{array}{cccc} \begin{array}{c} (A \rightarrow B) \\ A \\ \hline B \end{array} & \begin{array}{c} (A \rightarrow \neg B) \\ A \\ \hline \neg B \end{array} & \begin{array}{c} (\neg A \rightarrow B) \\ \neg A \\ \hline B \end{array} & \begin{array}{c} (\neg A \rightarrow \neg B) \\ \neg A \\ \hline \neg B \end{array} \\ \\ \begin{array}{c} (A \rightarrow B) \\ \neg B \\ \hline \neg A \end{array} & \begin{array}{c} (A \rightarrow \neg B) \\ B \\ \hline \neg A \end{array} & \begin{array}{c} (\neg A \rightarrow B) \\ \neg B \\ \hline A \end{array} & \begin{array}{c} (\neg A \rightarrow \neg B) \\ B \\ \hline A \end{array} \end{array}$$

Indirekter Beweis

Ein häufig verwendete Beweismöglichkeit ist der *indirekte Beweis* oder *Widerspruchsbeweis*, der auf den beiden Prinzipien der klassischen zweiwertigen Logik, dem ausgeschlossenen Widerspruch und dem ausgeschlossenen Dritten beruht.

Man möchte eine Aussage A beweisen und zeigt dazu, dass es nicht sein kann, dass die Negation von A , also $\neg A$ gilt. Wegen dem ausgeschlossenen Dritten muss dann A gelten. Um zu beweisen, dass $\neg A$ nicht gilt, zeigt man, dass aus $\neg A$ ein Widerspruch folgt: d. h. man findet

eine beliebige Aussage B , so dass aus A einerseits B folgt und andererseits auch $\neg B$ folgt. Dies macht man in der Regel so, dass man annimmt, dass $\neg A$ gilt (also dass A nicht gilt), und dann unter dieser Annahme zeigt, dass sowohl B gilt als auch B nicht gilt. Schematisch sieht der Widerspruchsbeweis also folgendermaßen aus:

Prämisse	$(\neg A \rightarrow B)$
Prämisse	$(\neg A \rightarrow \neg B)$
Konklusion	A

Ein klassisches Beispiel ist der Beweis der Irrationalität von $\sqrt{2}$:

Man möchte die Aussage A „ $\sqrt{2}$ ist irrational“ beweisen und nimmt dazu die Gültigkeit der Negation $\neg A$, also „ $\sqrt{2}$ ist rational“ an. Dies bedeutet, dass sich $\sqrt{2}$ als Bruch $\frac{p}{q}$ mit natürlichen Zahlen p und q schreiben lässt. Aus $\neg A$ folgt also die Aussage B : „Die Zahlen $p' := \frac{p}{\text{ggT}(p,q)}$ und $q' := \frac{q}{\text{ggT}(p,q)}$ sind teilerfremd“. Es gilt dann $\sqrt{2} = \frac{p'}{q'}$ und durch Quadrieren erhält man $p'^2 = 2q'^2$. Also ist p' gerade und damit p'^2 durch 4 teilbar, und daher ist auch q'^2 gerade und somit q' gerade. Unter der Annahme $\neg A$ folgt also auch die Gültigkeit von $\neg B$: „Die Zahlen p' und q' sind nicht teilerfremd“. Damit folgt aus $\neg A$ ein Widerspruch, und daher ist A bewiesen.

Manchmal gibt es eine Kurzform des indirekten Beweises, wenn man für die Aussage B die Aussage A selbst nehmen kann. Dann muss man natürlich nicht mehr zeigen, dass aus $\neg A$ auch $\neg A$ folgt, und der indirekte Beweis verkürzt sich zu folgendem Schema:

Prämisse	$(\neg A \rightarrow A)$
Konklusion	A

Wie das Auswahlaxiom hat auch der Beweis durch Widerspruch einen etwas mystischen Charakter. In schwächeren Logiken als der klassischen zweiwertigen Aussagenlogik funktioniert er auch nicht. Weitergehend anerkannt ist das Prinzip des ausgeschlossenen Widerspruchs; daraus erhält man die Korrektheit der folgenden Schlussweise:

Prämisse	$(\neg A \rightarrow B)$
Prämisse	$(\neg A \rightarrow \neg B)$
Konklusion	$\neg \neg A$

Aus dem Prinzip des ausgeschlossenen Dritten erhält man die Doppelnegationsregel, die es erlaubt von $\neg \neg A$ auf A zu schließen. In schwächeren Logiken (z. B. dem Intuitionismus) gilt dies nicht.

3.5 Quantoren und Variable

Wenn man auf der Ebene der Aussagenlogik bliebe, käme man mit dem Beweisen nicht besonders weit. Man muss sich auch die Struktur der Aussagen ansehen. Auch dazu gibt es Redeweisen, Symbole und als gültig angesehene Regeln (im Rahmen der bereits erwähnte Prädikatenlogik). Viele Aussagen behaupten etwas über Objekte, z. B. dass gewisse Eigenschaften gelten, und zwar in einer Art und Weise, die etwas über die Anzahl der Objekte aussagen, auf die die Behauptung zutrifft. Solche Formulierungen heißen *Quantifizierungen*. Einige der wichtigsten sind:

für alle natürlichen Zahlen gilt ... oder: für jede natürliche Zahl gilt ...
für keine natürlichen Zahl gilt ...
es gibt eine natürliche Zahl, so dass ... oder: für mindestens eine natürliche Zahl gilt ...
für endlich viele natürliche Zahlen gilt ...
für unendlich viele natürliche Zahlen gilt ...
für fast alle natürlichen Zahlen gilt ... (bedeutet in der Regel: für alle natürlichen Zahlen bis auf endlich viele gilt)

Dies waren nun Beispiele für *beschränkte Quantifizierungen*, bei denen die Art der betrachteten Objekte spezifiziert ist, im Beispiel stets die natürlichen Zahlen. Daneben gibt es auch *unbeschränkte Quantifizierungen* wie „es gibt etwas, so dass ...“ oder „für alle gilt ...“, die benutzt werden, wenn aus dem Kontext heraus klar ist, um welche Art von Objekten es sich handelt.

Für alle oben angegebenen Formulierungen gibt es natürlich sprachliche Varianten. Große Vorsicht ist in zweierlei Hinsicht geboten:

Die Verwendung des unbestimmten Artikels:

In der Mathematik bedeutet „es gibt ein ...“ stets „es gibt mindestens ein ...“. Dafür, dass es ein einziges gibt, verwendet man in der Mathematik die Wendungen „es gibt genau ein ...“ oder „es gibt ein und nur ein ...“.

genau ein

Der unbestimmte Artikel „ein“ ohne Zusatz wird oft auch wie eine Quantifizierung benutzt und zwar meistens als universelle. Oft taucht diese Art der Quantifizierungen in Formulierungen wie „Sei f eine Abbildung“ auf. Dies bedeutet: „Für alle Abbildungen gilt (und solch eine Abbildung sei im Folgenden f genannt)“.

Oder „Eine Menge von reellen Zahlen heißt beschränkt genau dann, wenn ...“ bedeutet: „Für jede Menge von reellen Zahlen gilt: sie heißt beschränkt genau dann, wenn ...“

Manchmal kann der unbestimmte Artikel aber auch für eine partikuläre Quantifizierung, also ein „es gibt ein ...“ stehen. Zum Beispiel in dem Satz: „Eine Folge reeller Zahlen konvergiert gegen eine reelle Zahl genau dann, wenn sie eine Cauchy-Folge ist.“ Dies bedeutet: „Für jede Folge reeller Zahlen gilt: Es gibt genau dann eine reelle Zahl, gegen die die Folge konvergiert, wenn ...“. Der erste unbestimmte Artikel „eine“ ist also ein universelle Quantifizierung, der zweite eine partikuläre.

Außerdem kann man den unbestimmten Artikel manchmal schwer von dem Zahlwort „ein“ unterscheiden. Gesprochen unterscheidet es sich durch die Betonung. Zum Beispiel bedeutet „Eine Primzahl ist gerade“ soviel wie: „Es gibt genau eine Primzahl, die gerade ist“. Dagegen bedeutet „Eine Quadratzahl ist ungerade oder durch 4 teilbar“ soviel wie „Jede Quadratzahl ist ungerade oder durch 4 teilbar“.

Wegen der möglichen Unklarheit sollte man den unbestimmten Artikel als Quantifizierung vermeiden und eine der am Anfang angegebenen eindeutigen Formulierungen verwenden!

Verschachtelte Quantifizierungen:

Bei mehreren Quantifizierungen kommt es auf die Reihenfolge an. Zum Beispiel gilt der Satz „Für jeden Menschen gibt es einen Schuh, der ihm passt“, aber nicht der Satz „Es gibt einen Schuh, der für jeden Menschen passt“. In der Umgangssprache wird die Reihenfolge nicht immer klar (z. B. „Ein Schuh passt jedem Menschen“). Dies ist tatsächlich auch immer wieder in der Mathematik ein Problem, weil Quantifizierungen zum Teil gar nicht, zum Teil vor und zum Teil hinter eine Aussage geschrieben werden.

Für zwei Quantifizierungen gibt es viel benutzte Symbole, nämlich:

Existenzquantor $\exists$ und *Allquantor* $\forall$

$\exists \forall$

Beide werden immer nur in Verbindung mit einer Variablen benutzt, über die in der auf den Quantor folgenden Formel etwas ausgesagt wird. Bei beschränkter Quantifizierung steht hinter der Variablen, aus welcher Menge die betrachteten Objekte stammen (was auch *relativierter Quantor* genannt wird).

$$\exists x \in \mathbb{N} (x \text{ ist Primzahl und } x > 1000)$$

heißt: „Es gibt eine natürliche Zahl, die eine Primzahl größer als 1000 ist.“

$$\forall y \in \mathbb{N} \ y \neq y + 1$$

heißt: „Jede natürliche Zahl ist verschieden von ihrem Nachfolger.“

Wenn man die Quantoren vor die Aussage schreibt, auf die sie sich beziehen, wird auch die Reihenfolge der Quantifizierungen klar.

$$\forall x \in \mathbb{N} \exists y \in \mathbb{N} x < y$$

ist eine wahre Aussage, weil man zu jeder natürlichen Zahl eine größere natürliche Zahl findet (z. B. zu n den Nachfolger $n + 1$).

$$\exists y \in \mathbb{N} \forall x \in \mathbb{N} x < y$$

gilt dagegen nicht, weil es keine größte natürliche Zahl gibt.

Geschachtelte Quantifizierungen implizieren nicht, dass die quantifizierten Objekte verschieden sind:

$$\forall x \exists y x \neq y$$

sagt aus, dass es zu jedem Ding ein dazu verschiedenes Ding gibt, also dass es mindestens zwei Dinge gibt. Dagegen behauptet

$$\exists x \forall y x \neq y$$

die Existenz eines Dinges, das von allem verschieden ist, also auch von sich selbst (was offenbar falsch ist).

Dies waren zudem Beispiele für unbeschränkte Quantifizierungen. Worauf sich die unbeschränkten Quantoren beziehen, muss vorher festgelegt werden oder aus dem Kontext klar sein. Ob die durch die Formel $\forall x \exists y x \neq y$ gemachte Aussage gilt oder nicht, hängt davon ab, auf welche Situation sie sich bezieht.

Im mathematischen Alltag werden Quantoren manchmal auch hinter die Aussage geschrieben, auf die sie sich beziehen (was bei geschachtelten Quantifizierungen ungeschickt ist). Manche fügen Zusatzzeichen ein, um die Struktur einer quantifizierten Aussage deutlicher hervortreten zu lassen: etwa Klammern um relativierte Quantoren oder Doppelpunkte danach, also etwa $(\forall x \in \mathbb{N}) x < x + 1$ oder $\forall x \in \mathbb{N} : x < x + 1$.

Es gibt noch eine andere verbreitete Symbolschreibweise für einen Quantor, nämlich $\exists!$ für „es gibt genau ein ...“, z. B. $\exists! x \in \mathbb{N} x^2 = 4$.

$\exists!$

Vor allem in der Analysis gibt es noch eine verkürzte Schreibweise für relativierte Quantoren, nämlich z. B. $\forall \varepsilon > 0$. Dies steht für $\forall \varepsilon \in \{r \in \mathbb{R} \mid r > 0\}$. Hier ist gewissermaßen aus dem Kontext „Analysis“ klar, dass die Variablen sich auf die reellen Zahlen beziehen, und durch die Größenangaben wird der Bereich nur noch weiter eingeschränkt. Die analogen Schreibweisen existieren für die anderen Quantoren und Ordnungssymbole, also etwa $\exists! y \leq -2 \quad y^2 = 4$.

Auch für natürliche Zahlen werden solche abkürzenden Schreibweisen benutzt, z. B. bei der Definition der Konvergenz einer Folge $(a_i)_{i \in \mathbb{N}}$ gegen den Grenzwert c schreibt man gerne:

$$\forall \varepsilon > 0 \exists n \in \mathbb{N} \forall m > n \quad |a_m - c| < \varepsilon$$

Hier ergibt sich aus dem Kontext, dass mit m eine natürliche Zahl gemeint sein muss und der Quantor „ $\forall m > n$ “ als „ $\forall m \in \{x \in \mathbb{N} \mid x > n\}$ “ zu lesen ist, da m im auf die Quantoren folgenden Ausdruck als Index eines Folgenglieds auftaucht, wofür durch die Angabe $(a_i)_{i \in \mathbb{N}}$ hier nur natürliche Zahlen zugelassen sind.

Regeln für Quantoren

Für den Umgang mit Quantoren gelten nun ebenfalls Regeln. Da es schwierig ist, diese Regeln in Umgangssprache zu formulieren, werde ich sie hier in formaler Schreibweise wiedergeben, ohne allerdings den Formalismus in seinen Einzelheiten zu erklären. Nur soviel dazu: Im Folgenden sei $\varphi(x)$ ein Ausdruck, in dem eine (einzige) *freie Variable* x vorkommt, z. B. „ x ist gerade“. Durch *Abquantifizieren von x* , also durch Voranstellen eines Quantors „es gibt ein x in M “ oder „für alle x in M “ soll zudem $\varphi(x)$ zu einer Aussage $\exists x \in M \varphi(x)$ bzw. $\forall x \in M \varphi(x)$ werden.

Außerdem kann man für die Variable x Namen von Elementen einsetzen; man kann also z. B. $\varphi(m)$ dafür schreiben, dass die durch φ beschriebene Eigenschaft auf m zutrifft. Sowohl $\exists x \varphi(x)$, als auch $\forall x \varphi(x)$ und $\varphi(m)$ sind Aussagen, d. h. im betrachteten Kontext sind sie wahr oder falsch. Dagegen ist „ $\varphi(x)$ “ im allgemeinen keine Aussage; man kann nicht sagen, dass $\varphi(x)$ stimmt oder nicht stimmt, da sich x auf kein bestimmtes Objekt bezieht. Solange man nicht weiß, wofür x stehen soll, hat die Behauptung „ x ist gerade“ keinen Sinn. Damit solch ein Ausdruck einen Sinn bekommt, muss x abquantifiziert sein oder der Name für ein festes Objekt sein.

Zunächst die Regeln, die sich aus der Definition der Quantoren ergeben:

- Die Aussage $\exists x \in M \varphi(x)$ gilt genau dann, wenn es ein Element m in der Menge M gibt, so dass $\varphi(m)$ gilt.
- Die Aussage $\forall x \in M \varphi(x)$ gilt genau dann, wenn $\varphi(m)$ für jedes einzelne Element m in der Menge M gilt.

Konkret hat dies aber die folgenden Auswirkungen:

- Wenn man eine Existenzaussage $\exists x \in M \varphi(x)$ bewiesen hat, dann kann man einem der (unbekannten) Elemente, deren Existenz behauptet wird, einen Namen m geben. Dieser Name muss „neu“ sein, d. h. er darf nicht schon vorher für ein spezielles Element von M verwendet worden sein.
- Wenn man eine Aussage $\varphi(m)$ für ein beliebiges Element m von M bewiesen hat, ohne dass irgendwo festgelegte besondere Eigenschaften von m eingegangen sind, dann hat man $\forall x \in M \varphi(x)$ bewiesen.

Tabelle 4 gibt einige wichtige Regeln für Quantoren an.

Tabelle 4: PRÄDIKATENLOGISCHE GESETZE

$\varphi(x, y)$ bedeutet, dass φ ein mathematischer Ausdruck ist, in dem sowohl x als auch y freie Variable sind, und der durch „Abquantifizieren“ von x und y zu einer Aussage wird.	
Abschwächung:	Falls $M \neq \emptyset$, so gilt $\forall x \in M \psi(x) \implies \exists x \in M \psi(x)$
Vertauschungsregeln	$\forall x \in M \forall y \in N \varphi(x, y) \iff \forall y \in N \forall x \in M \varphi(x, y)$ $\exists x \in M \exists y \in N \varphi(x, y) \iff \exists y \in N \exists x \in M \varphi(x, y)$ $\exists x \in M \forall y \in N \varphi(x, y) \implies \forall y \in N \exists x \in M \varphi(x, y)$ die umgekehrte Implikation gilt i. Allg. nicht!!
Regeln von de Morgan	$\neg (\exists x \in M \psi(x)) \iff \forall x \in M \neg \psi(x)$ $\neg (\forall x \in M \psi(x)) \iff \exists x \in M \neg \psi(x)$
Distributivgesetze	$\forall x \in M (\psi(x) \wedge \theta(x)) \iff (\forall y \in M \psi(y) \wedge \forall z \in M \theta(z))$ $(\exists x \in M \psi(x) \vee \exists y \in M \theta(y)) \iff \exists z \in M (\psi(z) \vee \theta(z))$ $\exists x \in M (\psi(x) \wedge \theta(x)) \implies (\exists y \in M \psi(y) \wedge \exists z \in M \theta(z))$ $(\forall x \in M \psi(x) \vee \forall y \in M \theta(y)) \implies \forall z \in M (\psi(z) \vee \theta(z))$ die umgekehrten Implikationen gelten i. Allg. nicht!!

Man kann übrigens Quantoren vor Ausdrücke schreiben, in denen die Quantifizierungsvariable gar nicht vorkommt, also etwa „für alle x gilt, dass 4 eine gerade Zahl ist“. Solche „sinnlosen“ Quantifizierungen bewirken nichts:

Falls x in φ nicht vorkommt, so gilt $\forall x \varphi \iff \exists x \varphi \iff \varphi$.

Die konkreten Namen x und y in den Regeln oben sind gewissermaßen „Variablen für Variablen“; sie haben keine eigene Bedeutung. Man könnte ebenso gut ein Kästchen schreiben. Es gilt also $\forall x \varphi(x) \iff \forall y \varphi(y)$ und $\exists x \varphi(x) \iff \exists y \varphi(y)$; und alle Regeln gelten ebenso, wenn man Variablen mit anderen Namen für x oder y einsetzt. Wichtig ist nur, dass man eindeutig erkennt, welcher Quantor sich auf was bezieht. Bei mehreren und insbesondere bei geschachtelten Quantifizierungen sollte man am besten für jeden Quantor eine eigene Variable benutzen, und diese Variable sollte am besten im Kontext nicht in anderer Bedeutung vorkommen. Es ist zwar genau definiert, was der Wirkungsbereich eines Quantors ist, und außerhalb des Wirkungsbereichs ist ein Variablenname für andere Benutzung frei; Doppelbenutzungen können aber leicht zu Verwirrung

Nur darf man nicht (oder nur unter besonderen Umständen) einen Variablennamen durch einen bereits vorkommenden ersetzen. Wenn man in $\forall x \exists y (\varphi(x) \wedge \psi(y))$ die Variable y durch x ersetzen wollte, würde sich der Sinn ändern. Man kann sie aber problemlos durch eine „neue“ Variable ersetzen, z. B. z .

☞ Gerne verwendet man bestimmte Buchstaben für Variablen eines bestimmten Typs. Zum Beispiel nennt man Variablen für natürliche Zahlen gerne m und n , Variablen für Funktionen gerne f, g, h . Dies ist nichts, was selbstverständlich wäre. Man kann nicht, weil man den Buchstaben n verwendet, schon daraus schließen, dass n für eine natürliche Zahl steht. Wenn es um Mengen X, Y, Z geht und man x als Variable für Elemente von X und y als Variable für Elemente von Y nimmt, ist damit noch nicht gesagt, dass z automatisch als Variable für Elemente aus Z steht. Bei jeder Benutzung einer Variablen muss irgendwo dazugesagt sein, wofür sie steht.

Konstanten und Variablen

In der mathematischen Formelsprache werden Namen für Objekte verwandt. Manche davon sind *Konstanten*, d. h. sie stehen stets für das gleiche Objekt. Zum Beispiel π für die Kreiszahl, $\mathbb{N}$ für die Menge der natürlichen Zahlen, „sin“ für die Sinusfunktion oder $=$ für die Gleichheitsrelation. Es kann aber sein, dass man sich zum Beispiel in einer Situation befindet, in der die Kreiszahl nicht auftritt. Dann wird π durchaus auch für anderes verwendet. Die wissenschaftlichen Typografiereregeln sehen vor, dass Konstanten in geraden Buchstaben geschrieben werden sollten.

Andere Namen stehen nicht immer für das gleiche Objekt und werden *Variable* genannt (und üblicherweise mit geneigten Buchstaben geschrieben). Obwohl kein prinzipieller Unterschied vorliegt, kann man doch subjektiv zwei verschiedene Gebrauchsweisen der Variablen unterscheiden. In der einen Weise haben die Variablen tendenziell einen kurzen Wirkungsbereich und sind explizit abquantifiziert (z. B. in Sätzen wie „es gibt eine Zahl n mit der Eigenschaft ...“ oder als Indexvariablen in Formeln wie $\sum_{i=1}^n \dots$). In der anderen Gebrauchsweise gibt man einem beliebigen (also variablen) Objekt eine gewissen Art für einen längeren Zeitraum einen festen Namen, z. B. durch die Formulierung „Sei M eine Menge“, gerne mit dem Zusatz „beliebig, aber fest“. Man behandelt M dann, als wäre es eine Konstante, in Wirklichkeit befindet man sich aber im Bereich eines umfassenden Allquantors „für alle M “.

Wenn Ausdrücke mit nicht explizit abquantifizierten Variablen als Aussagen verwendet werden, gilt die Konvention, dass die Variablen universell abquantifiziert sein sollen. Wenn ich also z. B. für Aussagen A schreibe: $A \iff \neg\neg A$, dann ist damit gemeint, dass diese Regel für alle Aussagen A gilt. Das in der ein oder anderen Weise dabeistehende oder dazugedachte „sei A eine beliebige Aussage“ beinhaltet den Allquantor.

3.6 Verneinungen

Vor allem um die Beweistechnik des indirekten Beweises anwenden zu können, muss man mathematische Aussagen verneinen können. Natürlich kann man dies immer durch Voranstellen der Phrase „es ist nicht der Fall, dass ...“ bzw. durch Voranstellen eines Negationszeichens erreichen. Dies hilft aber in der Regel nicht für den Beweis, man muss das Negationszeichen

anschließend noch „nach innen ziehen“. Dazu dienen die Doppelnegationsregel, die Regeln von de Morgan (für Junktoren und für Quantoren) und die Definition der Implikation. Hier sind nochmals alle relevanten Regeln zusammengefasst:

$\neg\neg A$	wird zu	A
$\neg(A \wedge B)$	wird zu	$(\neg A \vee \neg B)$
$\neg(A \vee B)$	wird zu	$(\neg A \wedge \neg B)$
$\neg(A \rightarrow B)$	wird zu	$(A \wedge \neg B)$
$\neg(A \leftrightarrow B)$	wird zu	$((A \wedge \neg B) \vee (A \wedge \neg B))$
$\neg(\exists x \in M \varphi(x))$	wird zu	$\forall x \in M \neg\varphi(x)$
$\neg(\forall x \in M \psi(x))$	wird zu	$\exists x \in M \neg\psi(x)$

Beispiel 1:

Nehmen wir an, dass n eine natürliche Zahl ist, und betrachten wir die Aussage:

Wenn n eine Quadratzahl ist oder durch 5 teilbar ist, dann lässt n bei Division durch 5 den Rest 0 oder 1 oder 4.

Diese Aussage hat die logische Form $((A \vee B) \rightarrow (C \vee D \vee E))$ mit A für „ n ist eine Quadratzahl“, B für „ n ist durch 5 teilbar“, C, D, E für „ n lässt sie bei Division durch 5 den Rest 0“ bzw. „den Rest 1“ bzw. „den Rest 4“.

Die Verneinung lässt sich dann mit den angegebenen Regeln Schritt für Schritt formal umformen:

$$\begin{aligned} & \neg((A \vee B) \rightarrow (C \vee D \vee E)) \\ & (A \vee B) \wedge \neg(C \vee D \vee E) \\ & (A \vee B) \wedge \neg C \wedge \neg D \wedge \neg E \end{aligned}$$

Sprachlich ergibt sich also als Verneinung der obigen Aussage: „ n ist eine Quadratzahl oder durch 5 teilbar, und n lässt bei Division durch 5 nicht den Rest 0 und nicht den Rest 1 und nicht den Rest 4.“ Wenn man nun noch weiß, dass als Reste sowieso nur 0, 1, 2, 3, 4 in Frage kommen, kann man die Aussage weiter vereinfachen zu: „ n ist eine Quadratzahl oder durch 5 teilbar, und n lässt bei Division durch 5 den Rest 2 oder den Rest 3.“

Mit etwas Übung kann man solche Verneinungen natürlich direkt auf der sprachlichen Ebene durchführen; allerdings kann es bei komplizierten Aussagen helfen, sich die logische Struktur der Aussage klar zu machen.

Beispiel 2:

Das weiter oben angegebene Beispiel

$$\forall \varepsilon > 0 \exists n \in \mathbb{N} \forall m > n \quad |a_m - c| < \varepsilon$$

kann man ebenfalls mit den Regeln Schritt für Schritt verneinen zu:

$$\begin{aligned} & \neg(\forall \varepsilon > 0 \exists n \in \mathbb{N} \forall m > n \quad |a_m - c| < \varepsilon) \\ & \exists \varepsilon > 0 \neg \exists n \in \mathbb{N} \forall m > n \quad |a_m - c| < \varepsilon \\ & \exists \varepsilon > 0 \forall n \in \mathbb{N} \neg \forall m > n \quad |a_m - c| < \varepsilon \\ & \exists \varepsilon > 0 \forall n \in \mathbb{N} \exists m > n \quad \neg |a_m - c| < \varepsilon \\ & \exists \varepsilon > 0 \forall n \in \mathbb{N} \exists m > n \quad |a_m - c| \geq \varepsilon \end{aligned}$$

(Bei solchen Aussagen mit vielen Quantoren ist die Formelschreibweise sehr nützlich. Wenn man versuchen wollte, solch eine Aussage in natürlicher Sprache wiederzugeben ohne ungenau zu werden, kann man eigentlich nichts anderes tun als die Formel „vorzulesen“.)

4 Mathematische Terminologie

4.1 Satz, Definition, Beweis ...

In einem vielfach gepflegten Stil werden Mathematikvorlesungen, Lehrbücher und mathematische Abhandlungen in kleine, oft durchnummerierte Einheiten zerteilt. Die üblichsten sind:

Satz oder **Proposition** für ein mathematisches Ergebnis mit eigenständigem Interesse. Besonders wichtige Ergebnisse werden auch mit **Hauptsatz** oder **Theorem** bezeichnet (wobei die Abgrenzung oft unscharf ist und manche Autoren diesen Unterschied auch gar nicht machen). Wichtige Sätze tragen oft Namen („Hauptsatz der Differential- und Integralrechnung“, „Hilberts Nullstellensatz“, „Satz von Fubini“ ...). Bei ganz großen Mathematikern (Euler, Gauß, ...) ist die Zuordnung allerdings nicht immer eindeutig.

Als **Lemma** oder **Hilfssatz** werden Zwischenergebnisse bezeichnet, die in mehreren Beweisen gebraucht werden oder der Übersichtlichkeit halber aus größeren Beweisen ausgegliedert werden. Daneben gibt es auch Lemmata mit Namen (Zorn'sches Lemma, Hensels Lemma, Lemma von Gauß ...). Dies sind zumeist Ergebnisse, die sich für die Entwicklung einer ganzen Theorie als grundlegend erwiesen haben (der Name „Lemma“ ist dann als Ehrenbezeichnung zu verstehen).

Als **Folgerung** oder **Korollar** aus einem Satz wird ein Ergebnis bezeichnet, welches keinen eigenständigen Beweis mehr braucht, sondern dessen Gültigkeit sofort oder fast sofort aus dem vorher bewiesenen Satz eingesehen wird.

Die Verwendung von **Vermutung** (engl: **conjecture**), **Beispiel** und **Bemerkung** sollte klar sein.

Eine exakte **Definition** führt einen neuen Begriff oder ein neues Symbol so ein, dass man überall da, wo der Begriff oder das Symbol auftaucht, ihn oder es durch die definierende Beschreibung ersetzen könnte, ohne dass sich die Aussage inhaltlich ändert. Bei mathematischen Objekten werden die definierenden Eigenschaften auch gerne **Axiome** genannt.

Wenn ein einzelnes Objekt definiert wird (z. B. die Zahl 5 als Nachfolger der Zahl 4), dann bedeutet dies, dass einem bereits benennbaren Objekt ein neuer Name gegeben wird. Nach solch einer Definition gilt also eine Gleichheit zwischen dem durch den neuen Namen benannten Objekt und dem durch den alten Namen bezeichneten Objekt. Manchmal wird eine solche Definition in Formelschreibweise einfach durch diese Gleichung wiedergegeben; oft wird irgendwie auf den definitorischen Charakter hingewiesen, z. B. durch einen Doppelpunkt auf der zu definierenden Seite, also $5 := 4 + 1$ oder $4 + 1 =: 5$ (dies ist eine aus der Informatik übernommene, also recht neue Schreibweise). Andere verbreitete Versionen sind $5 =_{\text{df}} 4 + 1$ und $5 \equiv 4 + 1$.

Wenn eine Relation definiert wird, wird dadurch eine Äquivalenz aufgestellt, z. B. zwischen der Aussage A „die natürliche Zahl m teilt die natürliche Zahl n “ und der Aussage B „ m und n sind natürliche Zahlen und es gibt eine natürliche Zahl k mit $k \cdot m = n$ “. Hierfür gibt es die analoge Schreibweise der „definierenden Äquivalenz“ $A : \Longleftrightarrow B$. Wenn klar ist, dass es sich um eine Definition handelt, spart man sich umgangssprachlich gerne das „genau dann“ und sagt z. B. „eine natürliche Zahl m teilt eine natürliche Zahl n , wenn es eine natürliche Zahl k mit $k \cdot m = n$ gibt“ und versteht dies so, dass solch eine Definition immer umfassend sein soll, dass es also keine anderen Umstände gibt, unter denen eine natürliche Zahl eine andere teilt.

Daneben gibt es eine Art „beschreibende Definition“ (in der nicht-axiomatischen Mathematik), welche Eigenschaften von Objekten nahelegt, ohne sie genau festzulegen, und die keine exakte Definition ist. So z. B. „Ein Punkt ist, was keine Teile hat“ bei Euklid oder „Unter einer ‚Menge‘ verstehen wir jede Zusammenfassung M von bestimmten wohlunterschiedenen Objekten unserer Anschauung oder unseres Denkens (welche die ‚Elemente‘ von M genannt werden) zu einem Ganzen“ bei Cantor.

Ein mathematischer Satz wird erst dann als gültig angesehen, wenn er bewiesen ist. In der Theorie gibt es eine exakte Definition davon, was ein **Beweis** ist, nämlich eine Abfolge von mathematischen Aussagen, die

- entweder Axiome sind
- oder Aussagen sind, die bereits bewiesen sind oder als bewiesen geltend

- oder Aussagen sind, die aus vorherigen Aussagen in der Beweiskette durch einen korrekten logischen Schluss folgen. Hierfür dürfen nur gewisse, fest vorgegebene korrekte Schlussweisen verwendet werden.

In der Praxis kann man aber nur ganz einfache Beweise wirklich in dieser Form ausführen; zu komplizierteren Sätzen käme man aus Zeitgründen nicht und weil diese Art des Beweisens todlangweilig wäre. Stattdessen erlaubt man sich „größere“ Schritte, die aber nachvollziehbar sein müssen.

Was nun als „nachvollziehbar“ angesehen wird, hängt sehr von der Situation ab. In der Praxis gibt es somit keine klaren Kriterien dafür, was ein gültiger Beweis ist. Man könnte vielleicht sagen, dass ein akzeptierter Beweis dann vorliegt, wenn die relevanten Personen davon überzeugt sind, dass man ihn in die oben beschriebene theoretische Form eines Beweises bringen könnte. Ob ein Beweis akzeptiert wird oder nicht, hängt also davon ab, wer ihn vorlegt und wer ihn beurteilt. Insbesondere wird man von einem Anfänger mehr und ausführlichere Zwischenschritte verlangen als von einem etablierten Mathematiker; ebenso wird man in einer Vorlesung anders beweisen als in einem Fachartikel.

Einen Beweis muss man außerdem verteidigen können: Wenn jemand einen Beweisschritt anzweifelt (und sei es nur aus didaktischen Gründen), dann muss man in der Lage sein, immer feinere Zwischenschritte einfügen zu können, letztendlich bis auf die Ebene des theoretischen Beweisbegriffs hinab.

Natürlich kann es vorkommen, dass ein Beweis eine Zeitlang als gültig angesehen wird und dass dann doch noch eine irreparable Lücke entdeckt wird. Ob ein Satz gilt oder nicht, hängt nicht von unserem Wissen ab, und auch, ob ein Beweis richtig ist oder falsch, wird üblicherweise als eine absolute Eigenschaft angesehen. Ob aber ein Resultat mit Beweis von der Gemeinschaft der Mathematiker/innen als gültig angesehen wird oder nicht, kann sich im Laufe der Zeit ändern.

4.2 Weiterer mathematischer Sprachgebrauch

Noch ein paar Bemerkungen zu häufig gebrauchten Wörtern, Formulierungen und Schreibweisen:

genau: Wird verwendet für „nicht mehr und nicht weniger“. Zum Beispiel bedeutet „Die Menge M enthält genau die ganzen Zahlen“, dass alle ganzen Zahlen Elemente der Menge M sind und dass es keine anderen Elemente gibt. Der Satz „Die Menge M enthält die ganzen Zahlen“ lässt auch die Interpretation zu, dass es noch andere Zahlen in M gibt (es sei denn, er ist als Definition der Menge M gemeint). Ebenso heißt „genau ein“ soviel wie „mindestens ein und höchstens ein“.

Klammern: Klammern haben in der Mathematik zwei Rollen: Zum einen bilden sie festgefügte Symbole wie bei Tupeln oder Binomialkoeffizienten. Hierbei muss die Art der Klammern beachtet werden (rund, eckig, spitz, geschweift): Andere Klammern haben eine andere (oder keine) Bedeutung.

Zum andern werden Klammern für die Strukturierung von Formeln eingesetzt. Hier gibt es notwendige Klammern (z. B. bei nicht vertauschenden Rechenoperationen) und nicht-notwendige Klammern, die nur der besseren Lesbarkeit willen eingesetzt werden. Vor allem letztere werden oft einfach benutzt, ohne dass über die Verwendungsweise gesprochen wird. Für solche Strukturierungsklammern werden die verschiedenen Klammerarten ohne Unterschied verwendet. Auch die Klammergröße spielt keine Rolle.

oder: immer einschließend gebraucht, wenn es nicht „entweder ... oder ...“ heißt.

paarweise verschieden wird oft der Genauigkeit halber benutzt. „Die Elemente $a_1, \dots, a_n$ sind paarweise verschieden“ bedeutet, dass für je zwei verschiedene Indizes i und j auch die Elemente a_i und a_j verschieden sind, dass es sich also insgesamt um n verschiedene Objekte handelt. „Die Elemente $a_1, \dots, a_n$ sind verschieden“ kann man auch dahingehend verstehen, dass sie nicht alle gleich sind.

trivial: Dass ein Beweis trivial sei oder eine Aussage trivialerweise richtig, bedeutet je nach

Sprecher etwas im Bereich von „es ist offensichtlich“ über „es ist einfach“ bis „man braucht im Beweis keine neuen Ideen“. Daneben wird „trivial“ in der Mathematik auch in Fachbegriffen verwendet, z. B. „die triviale Gruppe“.

Zur Herkunft von „trivial“: In der Spätantike und im frühen Mittelalter bestand die höhere Ausbildung aus den sogenannten Sieben freien Künsten, den *Artes Liberales*: zunächst gewissermaßen als Grundstudium die drei sprachlich-logischen Disziplinen Grammatik, Rhetorik und Logik, die zusammen das *Trivium* („Dreiweg“) bildeten, dann als „Hauptstudium“ das *Quadrivium* („Vierweg“) der vier mathematischen Disziplinen Arithmetik, Geometrie, Harmonielehre und Astronomie. „Trivial“ war also, was aus dem Trivium, dem Grundstudium, bekannt war. Als im Hochmittelalter die Universitäten mit zunächst den Fakultäten Theologie, Jura und Medizin entstanden, wurden die Freien Künste einerseits zu einem Propädeutikum des Studiums, andererseits zum Vorläufer der Philosophischen Fakultäten, weswegen die geisteswissenschaftlichen Abschlüsse im angelsächsischen Bereich und seit der Bologna-Reform auch bei uns *Bachelor of Arts* und *Master of Arts* heißen.

Zu guter Letzt: Die Mathematik ist in ihrer Schreib- und Sprechweise nicht immer so exakt, wie man glauben könnte. Viele Sprech- und Schreibweisen werden benutzt, ohne dass sie immer im Detail erklärt werden. In der Regel sollte sich die Bedeutung aus dem Gebrauch ergeben, aber bei Unklarheiten sollte man stets nachfragen!

5 Anhang

5.1 Axiomatische Mengenlehre

Wenn man den Mengenbegriff zu weit fasst und „*volle Komprehension*“ zulässt, also für alle Eigenschaften E annimmt, dass es die Menge aller Objekte mit der Eigenschaft E gibt, so führt dies schnell zu Widersprüchen, z. B. zu der (zunächst von Zermelo entdeckten) Russell'schen Antinomie: Die dann existierende Menge $R := \{x \mid x \notin x\}$, also die Menge aller Mengen, die sich nicht selbst als Element enthalten, lässt sich mit der klassischen Aussagenlogik nicht in Einklang bringen, da man der Aussage $R \in R$ keinen der beiden Wahrheitswerte zuordnen kann.

(Ein etwas konkretes Beispiel für diese Problematik: In den meisten Büchern wird das Inhaltsverzeichnis des Buches nicht im Inhaltsverzeichnis aufgeführt, es ist also ein Verzeichnis, das sich nicht selbst aufführt. Wenn Sie versuchen wollten, ein Verzeichnis aller Verzeichnisse zu erstellen, die sich nicht selbst aufführen, müssten Sie dann Ihr Verzeichnis in Ihr Verzeichnis aufnehmen oder nicht?)

In Systemen axiomatischer Mengenlehre versucht man daher durch Axiome genau zu beschreiben, welche Mengen man zulässt und welche nicht. Das verbreitetste System ist das von Zermelo, Skolem und Fränkel, *ZFC* („C“ für Choice, d. i. das Auswahlaxiom). Darin gibt es nur noch die *eingeschränkte Komprehension* oder Aussonderung: Wenn man schon eine Menge M hat und eine (in einer gewissen, genau festgelegten Sprache beschreibbare) Eigenschaft E , dann kann man die Menge der Elemente von M mit der Eigenschaft E bilden.

Nun spricht man als Mathematiker immer wieder von „allen Mengen“ (z. B. beständig in diesem Skript, wenn es heißt, dass für jede Menge etwas gilt) und stellt sich damit auch die Gesamtheit aller Mengen vor. Diese kann keine Menge sein, weil es mit der Aussonderungsregel zusammen zur Russell'schen Antinomie führen würde, ist aber kein in sich widersprüchliches Konzept. Dafür benutzt man dann den Ausdruck *Klasse*, d. h. man spricht zum Beispiel von der „Klasse aller Mengen“. Jede Menge ist eine Klasse, aber es gibt Klassen, die keine Mengen sind, weil sie gewissermaßen zu groß dafür sind. In der axiomatischen Mengenlehre von Gödel und Bernays wird auch der Klassenbegriff formalisiert. Darin gibt es dann die Klasse aller Mengen, die sich nicht selbst als Element enthalten. Da diese Klasse aber keine Menge ist, tritt für sie das Problem mit der Russell'schen Antinomie nicht auf. Das Problematische an der „Russell-Menge“ R ist die *imprädikative Definition*: Die Menge, die definiert wird, kommt in der Definition selbst vor. In der Klassenversion ist dies aufgelöst. (Aber nicht jede imprädikative Definition ist problematisch!)

5.2 Ein Exkurs über das Auswahlaxiom

Schwierig ist, wie öfters in der Mathematik, der Umgang mit unendlichen Mengen, da man diesen nicht aus der Anschauung gewinnen kann. Welche Verhaltensweisen endlicher Mengen darf man auf unendliche übertragen? Es braucht also Axiome, um Regeln für den Umgang mit der Unendlichkeit festzulegen. Eine solche Regel ist das Auswahlaxiom, das einen etwas mystischen Charakter erlangt hat. (Eine andere solche Regel ist das weiter vorne besprochene Induktionsprinzip.) Hier wie anderswo gibt es Meinungsunterschiede zwischen den Mathematikern, die möglichst freizügige Regeln anwenden, solange kein Widerspruch auftritt (tendenziell die Mehrheit), und jenen, die eher restriktive Regeln bevorzugen, um sicherzustellen, dass keine Widersprüche auftreten werden und die betriebene Mathematik sinnvoll ist (tendenziell eher die Minderheit).

Für endliche Indexmengen I gilt das Auswahlaxiom (d. h. es folgt aus den anderen, unstrittigen Axiomen der Mengenlehre). Für unendliche Indexmengen liefert es aber die Existenz von Abbildungen, die man in Regel nicht konstruieren oder konkreter beschreiben kann. Es wird daher von einer Minderheit restriktiver arbeitenden Mathematiker abgelehnt. Wenn man Mathematik als ein Spiel mit Symbolen nach gewissen Regeln betrachtet, dann ist das Auswahlaxiom nur eine von mehreren Spielregeln. Wenn man aber der Meinung ist, dass Mathematik eine in ir-

gendeinem Sinne tatsächlich existierende Ideenwelt beschreibt, dann stellt sich tatsächlich auch die Frage, ob das Auswahlaxiom gilt oder nicht.

5.3 Die Peano–Axiome

Warum stimmen die oben angegebenen Eigenschaften der natürlichen Zahlen? Warum ist z. B. die Addition der natürlichen Zahlen tatsächlich kommutativ?

Dies ist wieder einmal eine eher philosophische Frage, weil nicht klar ist, was es bedeuten soll, dass die Axiome „stimmen“ oder dass etwas „tatsächlich“ gilt. Dazu müssten die natürlichen Zahlen irgendwo in der Natur existieren und man müsste (zumindest theoretisch) nachschauen können, ob die Axiome mit dem Verhalten der natürlichen Zahlen in der Realität übereinstimmen. In einer gewissen Weise ist dies auch so, weil die natürlichen Zahlen auch dazu dienen, Alltagsphänomene zu modellieren. Für kleine natürliche Zahlen (solche, die im Alltag auftreten) ist die Kommutativität der Addition eine Alltagerfahrung (der Gesamtpreis eines Einkaufs hängt nicht von der Reihenfolge ab, in der die Waren auf das Band an der Supermarktkasse gelegt werden).

Substantiellere Herangehensweisen an diese Fragen bestehen darin, die natürlichen Zahlen in anderen mathematischen Strukturen zu modellieren oder einfachere und damit vielleicht unmittelbar einsichtige Axiome zu formulieren. Ersteres ist zum Beispiel in der axiomatischen Mengenlehre möglich. In ihr kann man die natürlichen Zahlen auf eine Weise wiederfinden, die das Alltagsverständnis der natürlichen Zahlen als Anzahlen endlicher Mengen modelliert, und kann die oben angegebenen Eigenschaften (bzw. die unten folgenden Peano–Axiome) beweisen, allerdings auf Grundlage von unbewiesenen Axiomen an das Verhalten der Mengen. Auch geometrische Interpretationen sind möglich, in denen z. B. die Kommutativität der Addition aus der (anschaulich einsichtigen) Invarianz der Anzahl einer Menge von Punkten unter einer Spiegelung folgt.

Um die zweite Herangehensweise hat Peano sich in seiner Axiomatik der natürlichen Zahlen bemüht, den Peano–Axiomen, die von der Zahl 0 und der Nachfolgeroperation $N(x)$ ausgeht. Die Axiome besagen:

1. 0 ist eine natürliche Zahl und jede natürliche Zahl hat einen eindeutigen Nachfolger in der Menge der natürlichen Zahlen.
2. 0 ist nicht Nachfolger einer natürlichen Zahl, aber jede andere natürliche Zahl ist Nachfolger einer eindeutigen natürlichen Zahl.
3. Es gilt das Induktionsprinzip.

Dann kann man induktiv die Addition, die Multiplikation und die Ordnung definieren durch:

$$\begin{aligned} n + 0 &:= n & n + N(m) &:= N(n + m) \\ n \cdot 0 &:= 0 & n \cdot N(m) &:= (n \cdot m) + n \\ n > 0 &:\iff n \neq 0 & n > N(m) &:\iff n > m \text{ und } n \neq N(m) \end{aligned}$$

und man kann beweisen, dass die oben aufgeführten Eigenschaften gelten (und außerdem, dass die natürlichen Zahlen dadurch bis auf Isomorphie eindeutig bestimmt sind).

Die natürlichen Zahlen nach von Neumann

Ganz kurz angedeutet sei noch, wie man die natürlichen Zahlen in der Mengenlehre wiederfinden kann auf eine von John von Neumann eingeführten Art und Weise:

Man definiert $0 := \emptyset$ und die Nachfolgeroperation $N(n) := n \cup \{n\}$. Damit ist dann induktiv:

$$\begin{aligned}1 &:= N(0) = \emptyset \cup \{\emptyset\} = \{\emptyset\} = \{0\} \\2 &:= N(1) = \{\emptyset\} \cup \{\{\emptyset\}\} = \{\emptyset, \{\emptyset\}\} = \{0, 1\} \\&\vdots \\n+1 &:= N(n) = n \cup \{n\} = \{0, \dots, n-1\} \cup \{n\} = \{0, \dots, n\}\end{aligned}$$

Eine natürliche Zahl ist hier also eine Menge, die genauso viele Elemente enthält, wie die Zahl selbst angibt. Und es ist eine besondere solche Menge; ihre Elemente sind nämlich gerade die natürlichen Zahlen, die kleiner sind als sie. Man kann sich die natürlichen Zahlen nach dieser Definition vorstellen als Repräsentanten der Äquivalenzklassen der „Gleichmächtigkeitsrelation“.

Mit Hilfe der Axiome der Mengenlehre kann man nun beweisen, dass die Peano-Axiome gelten, wobei man ein besonderes Axiom braucht, um sicherzustellen, dass die Menge der natürlichen Zahlen (die mengentheoretisch meist ω genannt wird) überhaupt existiert und das Induktionsprinzip erfüllt.