

1. Etwas Logik und Mengenlehre

Bevor wir mit dem eigentlichen Inhalt der Vorlesung beginnen, müssen wir in diesem Kapitel kurz die exakte mathematische Sprache beschreiben, in der wir unsere Ergebnisse formulieren werden: die der Logik und Mengenlehre. Zentral hierbei sind die Begriffe der *Aussage* (in der Logik) und der *Menge* (in der Mengenlehre).

Da wir es hier mit den ersten beiden Begriffen überhaupt zu tun haben, die in der Mathematik vorkommen, können wir sie natürlich nicht durch bereits bekannte Dinge definieren oder mit bereits bekannten Resultaten ihre Eigenschaften herleiten. Wir müssen sie daher (wie schon in der Einleitung erwähnt) axiomatisch voraussetzen. Wir müssen *voraussetzen*, dass es sinnvoll ist, über logische Aussagen und deren Wahrheit zu reden, dass Mengen überhaupt existieren, dass man Mengen vereinigen und schneiden kann, aus ihnen Elemente auswählen kann, und noch einiges mehr. Wenn ihr euch zum Beispiel auf den Standpunkt stellt, dass ihr nicht an die Existenz von Mengen glaubt, wird euch niemand widerlegen können. Allerdings zweifelt ihr damit dann auch die Existenz der gesamten Mathematik an, wie sie heutzutage betrieben wird – und aus der Tatsache, dass ihr in dieser Vorlesung sitzt, schließe ich einmal, dass das nicht der Fall ist.

Glücklicherweise sind die Dinge, die wir benötigen, jedoch allesamt anschaulich sofort einleuchtend. Ich möchte es euch (und mir) daher ersparen, an dieser Stelle eine vollständige und präzise axiomatische Formulierung der Logik und Mengenlehre hinzuschreiben, zumal das momentan sicher mehr verwirren als helfen würde und außerdem gerade im Bereich der Logik auch zu sehr in die Philosophie abdriften würde. Stattdessen wollen wir uns in diesem (für den Rest der Vorlesung sehr untypischen) ersten Kapitel damit begnügen, die für uns wichtigsten Prinzipien und Notationen sowie beliebte Fehlerquellen in verständlicher Sprache zu erklären, auch wenn ein paar Dinge (insbesondere die Begriffsfestlegung – „Definition“ möchte ich es eigentlich gar nicht nennen – einer Aussage und einer Menge) dadurch recht schwammig klingen werden. Außerdem werden wir in Beispielen zur besseren Verdeutlichung bereits hier die reellen Zahlen und ihre einfachsten Eigenschaften (die euch sicherlich bekannt sein werden) benutzen, auch wenn wir diese erst später formalisieren werden. Da es sicher niemanden von euch verwirren wird, werden wir auch die Schreibweise „ $x \in \mathbb{R}$ “ für „ x ist eine reelle Zahl“ schon verwenden, bevor sie in den Notationen 1.12 und 1.14 offiziell eingeführt wird. Ab Kapitel 2 werden wir dann mit dem strukturierten Aufbau der Grundlagen der Mathematik beginnen, also alle Definitionen mit bereits eingeführten Notationen präzise formulieren und alle Aussagen aus vorherigen exakt beweisen.

1.A Logik

Beginnen wir also mit der Logik. Unter einer **Aussage** verstehen wir (grob gesagt) ein sprachliches Gebilde, das entweder wahr oder falsch ist (wobei wir in der Mathematik natürlich letztlich daran interessiert sind, *wahre* Aussagen herzuleiten – wenn wir später mathematische Sätze formulieren, ist damit also stets gemeint, dass die dort gemachte Aussage wahr ist). Wichtig sind auch sprachliche Gebilde, in denen freie **Variablen**, also Platzhalter, vorkommen, und die erst beim Einsetzen von Werten für diese Variablen Aussagen liefern. Man bezeichnet sie als **Aussageformen**.

Beispiel 1.1.

- (a) $1 + 1 = 2$ ist eine wahre, $1 + 1 = 3$ eine falsche, und $1 + 1$ überhaupt keine Aussage.
- (b) $x + 1 = 2$ ist eine Aussageform, die beim Einsetzen von $x = 1$ in eine wahre, beim Einsetzen jeder anderen reellen Zahl in eine falsche Aussage übergeht.

Bemerkung 1.2. Als Variablen in Aussageformen kann man beliebige Symbole benutzen. Üblich sind neben den normalen lateinischen Klein- und Großbuchstaben auch die griechischen Buchstaben, die wir zur Erinnerung hier auflisten:

A α alpha	B β beta	Γ γ gamma	Δ δ delta	E ε epsilon	Z ζ zeta	H η eta	Θ ϑ theta
I ι iota	K κ kappa	Λ λ lambda	M μ my	N ν ny	Ξ ξ xi	O o omikron	Π π pi
P ρ rho	Σ σ sigma	T τ tau	Y υ ypsilon	Φ ϕ phi	X χ chi	Ψ ψ psi	Ω ω omega

Oft verziert man Buchstaben auch noch mit einem Symbol oder versieht sie mit einem Index, um neue Variablen zu erhalten: So sind z. B. $x, x', \bar{x}, x_1, x_2, \dots$ alles Symbole für verschiedene Variablen, die zunächst einmal nichts miteinander zu tun haben (aber möglichst für irgendwie miteinander zusammenhängende Objekte eingesetzt werden sollten, wenn man den Leser nicht vollends verwirren will).

Notation 1.3 (Zusammengesetzte Aussagen). Sind A und B Aussagen, so lassen sich daraus wie folgt neue bilden:

Symbol	Wahrheitstafel				Bedeutung
A	w	f	w	f	
B	w	w	f	f	
$\neg A$	f	w			nicht A
$A \wedge B$	w	f	f	f	A und B
$A \vee B$	w	w	w	f	A oder B (oder beides): „nicht-ausschließendes Oder“
$A \Leftrightarrow B$	w	f	f	w	A und B sind gleichbedeutend / <i>äquivalent</i> , bzw. A genau dann, wenn B
$A \Rightarrow B$	w	w	f	w	aus A folgt B , bzw. wenn A dann B

Die sogenannte **Wahrheitstafel** in den mittleren vier Spalten ist dabei die eigentliche Definition der neuen zusammengesetzten Aussagen. Sie gibt in Abhängigkeit der Wahrheit von A und B (in den ersten beiden Zeilen) an, ob die zusammengesetzte Aussage wahr oder falsch ist.

Bemerkenswert ist hierbei wohl nur die Folgerungsaussage $A \Rightarrow B$, die keine Aussage über die Richtigkeit von A oder B separat macht, sondern nur sagt, dass B wahr ist, wenn auch A es ist. Ist hingegen A falsch, so ist die Folgerungsaussage $A \Rightarrow B$ stets wahr („aus einer falschen Voraussetzung kann man alles folgern“). So ist z. B. $0 = 1 \Rightarrow 2 = 3$ eine wahre Aussage. Wie schon erwähnt wollen wir uns in der Mathematik aber natürlich in der Regel mit wahren Aussagen beschäftigen, und neue wahre Aussagen aus alten herleiten. Gerade in Beweisen ist die übliche Verwendung der Notation $A \Rightarrow B$ daher, dass A eine bereits als wahr erkannte Aussage ist, und wir damit nun schließen wollen, dass auch B wahr ist.

Bemerkung 1.4 (Beweise mit Wahrheitstafeln). Wollen wir kompliziertere zusammengesetzte Aussagen miteinander vergleichen, so können wir dies auch mit Hilfe von Wahrheitstafeln tun. So ist für zwei Aussagen A und B z. B.

$$A \Rightarrow B \quad \text{äquivalent zu} \quad (\neg A) \vee B,$$

denn wenn wir in der Wahrheitstafel

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

mit Hilfe der Definitionen von $\neg$ und $\vee$ aus Notation 1.3 zunächst $\neg A$ und dann $(\neg A) \vee B$ berechnen, sehen wir, dass das Ergebnis mit $A \Rightarrow B$ übereinstimmt. Nach der Bemerkung aus Notation 1.3 ist dies auch anschaulich klar: Die Folgerungsaussage $A \Rightarrow B$ ist ja genau dann wahr, wenn A falsch (also $\neg A$ wahr) ist, oder wenn B wahr ist (oder beides).

Genauso zeigt man die ebenfalls einleuchtende Aussage, dass

$$A \Leftrightarrow B \quad \text{äquivalent zu} \quad (A \Rightarrow B) \wedge (B \Rightarrow A)$$

ist – was auch die übliche Art ist, wie man eine Äquivalenz zeigt: Man zeigt separat die beiden Folgerungen $A \Rightarrow B$ und $B \Rightarrow A$.

Notation 1.5. Folgerungen („ $\Rightarrow$ “) und Äquivalenzen („ $\Leftrightarrow$ “) sind natürlich zwei verschiedene Dinge, die man nicht durcheinanderwerfen darf (auch wenn das in der Schule wahrscheinlich manchmal nicht so genau genommen wird). Es hat sich jedoch in der Mathematik eingebürgert, bei *Definitionen* von Begriffen durch eine äquivalente, definierende Eigenschaft die Sprechweise „wenn“ anstatt des eigentlich korrekten „genau dann, wenn“ zu verwenden: So würde man z. B. als Definition des Begriffs einer geraden Zahl hinschreiben

„Eine ganze Zahl x heißt gerade, wenn $\frac{x}{2}$ eine ganze Zahl ist“,

obwohl man genau genommen natürlich meint

„Eine ganze Zahl x heißt *genau dann* gerade, wenn $\frac{x}{2}$ eine ganze Zahl ist“.

Eine gewöhnliche Folgerungsaussage wie z. B. die wahre Aussage

„Wenn eine ganze Zahl x positiv ist, dann ist auch $x + 1$ positiv“

ist dagegen immer nur in einer Richtung zu verstehen; hier wird also nicht behauptet, dass mit $x + 1$ auch x immer positiv sein muss (was ja auch falsch wäre).

Notation 1.6 (Quantoren). Natürlich kann man nicht nur zwei, sondern auch mehrere Aussagen wie in Notation 1.3 mit „und“ bzw. „oder“ verknüpfen – also die neue Aussage konstruieren, dass *jede* bzw. *mindestens eine* der ursprünglichen Aussagen wahr ist. Am einfachsten notiert man dies mit einer Aussageform A , in der eine freie Variable x vorkommt. Wir schreiben dies dann auch als $A(x)$ und setzen

Symbol	Bedeutung
$\forall x: A(x)$	für alle x gilt $A(x)$ (also eine „und“-Verknüpfung aller Aussagen $A(x)$)
$\exists x: A(x)$	es gibt ein x mit $A(x)$ (also eine „oder“-Verknüpfung aller Aussagen $A(x)$)

Die beiden Symbole $\forall$ und $\exists$ bezeichnet man als **Quantoren**. Beachte, dass diese beiden Quantoren *nicht* miteinander vertauschbar sind: So besagt z. B. die Aussage

$$\forall x \in \mathbb{R} \exists y \in \mathbb{R}: y > x$$

„zu jeder reellen Zahl x gibt es eine Zahl y , die größer ist“ (was offensichtlich wahr ist), während die Umkehrung der beiden Quantoren die Aussage

$$\exists y \in \mathbb{R} \forall x \in \mathbb{R}: y > x$$

„es gibt eine reelle Zahl y , die größer als jede reelle Zahl x ist“ liefern würde (was ebenso offensichtlich falsch ist). Der Unterschied besteht einfach darin, dass im ersten Fall zuerst das x gewählt werden muss und dann ein y dazu existieren muss (das von x abhängen darf), während es im zweiten Fall *dasselbe* y für alle x sein müsste.

Bemerkung 1.7. Jede Aussage lässt sich natürlich auf viele Arten aufschreiben, sowohl als deutscher Satz als auch als mathematische Formel. Die gerade eben betrachtete Aussage könnte man z. B. auf die folgenden (absolut gleichwertigen) Arten aufschreiben:

- (a) $\forall x \in \mathbb{R} \exists y \in \mathbb{R}: y > x$.
- (b) Es sei $x \in \mathbb{R}$. Dann gibt es ein $y \in \mathbb{R}$ mit $y > x$.
- (c) Zu jeder reellen Zahl gibt es noch eine größere.

Welche Variante man beim Aufschreiben wählt, ist weitestgehend Geschmackssache. Die Formulierung einer Aussage als deutscher Satz hat den Vorteil, dass wir sie oft leichter verstehen können, weil wir die deutsche Sprache schon länger kennen als die mathematische. Wenn wir uns jedoch erst einmal an die mathematische Sprache gewöhnt haben, wird auch sie ihre Vorzüge bekommen:

Sie ist deutlich kürzer und besser logisch strukturiert. Wir werden im Folgenden beide Schreibweisen mischen und jeweils diejenige wählen, mit der unsere Aussagen (hoffentlich) am einfachsten verständlich werden.

Wenn wir mathematische Symbole verwenden, müssen wir diese aber auch stets in ihrer korrekten Notation und nicht als „Abkürzungen“ für deutsche Wörter verwenden: Man würde die Aussage „2 und 4 sind gerade Zahlen“ sicher niemals schreiben als „ $2 \wedge 4$ sind gerade Zahlen“, und analog genauso wenig „Es gilt $x^2 \geq 0$ für alle $x \in \mathbb{R}$ “ als „Es gilt $x^2 \geq 0 \forall x \in \mathbb{R}$ “.

Bemerkung 1.8 (Negationen). Es ist wichtig zu wissen, wie man von einer Aussage das Gegenteil, also die Negation bzw. „Verneinung“ bildet. Da hierbei oft Fehler gemacht werden, wollen wir die allgemeinen Regeln hierfür im Folgenden kurz auflisten. Dabei könnte man (a), (b) und (c) wieder schnell mit Wahrheitstabellen zeigen; die Regeln (d) und (e) sind wie in Notation 1.6 analog zu (b) und (c) für die Verknüpfung mehrerer Aussagen:

- (a) $\neg(\neg A) \Leftrightarrow A$: Ist es falsch, dass A falsch ist, so bedeutet dies genau, dass A wahr ist.
- (b) $\neg(A \wedge B) \Leftrightarrow (\neg A) \vee (\neg B)$: Das Gegenteil von „ A und B sind richtig“ ist „ A oder B ist falsch“.
- (c) $\neg(A \vee B) \Leftrightarrow (\neg A) \wedge (\neg B)$: Das Gegenteil von „ A oder B ist richtig“ ist „ A und B sind falsch“.
- (d) $\neg(\forall x: A(x)) \Leftrightarrow \exists x: \neg A(x)$: Das Gegenteil von „für alle x gilt $A(x)$ “ ist „es gibt ein x , für das $A(x)$ falsch ist“.
- (e) $\neg(\exists x: A(x)) \Leftrightarrow \forall x: \neg A(x)$: Das Gegenteil von „es gibt ein x , für das $A(x)$ gilt“ ist „für alle x ist $A(x)$ falsch“.

Man kann also sagen, dass eine Verneinung dazu führt, dass „und“ mit „oder“ sowie „für alle“ mit „es gibt“ vertauscht werden. So ist z. B. das Gegenteil der Aussage

„In Frankfurt haben *alle* Haushalte Strom *und* fließendes Wasser“

die Aussage

„In Frankfurt *gibt* es einen Haushalt, der keinen Strom *oder* kein fließendes Wasser hat“.

Beispiel 1.9.

- (a) Wollen wir eine Folgerung $A \Rightarrow B$ verneinen, so können wir sie zunächst mit Bemerkung 1.4 zu $(\neg A) \vee B$ umformen, und erhalten nach Bemerkung 1.8 als Umkehrung dann $A \wedge \neg B$. Dies ist auch anschaulich einleuchtend: Die Folgerungsaussage „wenn A dann B “ ist genau dann falsch, wenn die Voraussetzung A zwar gilt, die Behauptung B aber nicht. Wir sehen also:

Die Verneinung einer Folgerung $A \Rightarrow B$ ist $A \wedge \neg B$

(und nicht etwa $A \Rightarrow \neg B$, wie man vielleicht denken könnte).

- (b) Eine oft vorkommende Anwendung der Regeln für die Verneinung von Aussagen ist der sogenannte **Widerspruchsbeweis** bzw. Beweis durch **Kontraposition**. Nach Bemerkung 1.4 gesehen ist die Folgerung $A \Rightarrow B$ („aus A folgt B “) gleichbedeutend mit $(\neg A) \vee B$. Damit ist diese Aussage nach Bemerkung 1.8 (a) auch äquivalent zu $(\neg(\neg B)) \vee (\neg A)$, also zu $\neg B \Rightarrow \neg A$. Mit anderen Worten: Haben wir eine Schlussfolgerung $A \Rightarrow B$ zu beweisen, so können wir genauso gut $(\neg B) \Rightarrow (\neg A)$ zeigen, d. h. *wir können annehmen, dass die zu zeigende Aussage B falsch ist und dies dann zu einem Widerspruch führen bzw. zeigen, dass dann auch die Voraussetzung A falsch sein muss*.

Beispiel 1.10. Hier sind zwei Beispiele für die Anwendung der Prinzipien aus Bemerkung 1.8 und Beispiel 1.9 – und auch unsere ersten Beispiele dafür, wie man Beweise von Aussagen exakt aufschreiben kann.

- (a) Einen Beweis durch Widerspruch könnte man z. B. so aufschreiben:

Behauptung: Für alle $x \in \mathbb{R}$ gilt $2x + 1 > 0$ oder $2x - 1 < 0$.

Beweis: Angenommen, die Behauptung wäre falsch, d. h. (nach Bemerkung 1.8 (c) und (d)) es gäbe ein $x \in \mathbb{R}$ mit

$$2x + 1 \leq 0 \quad (1) \quad \text{und} \quad 2x - 1 \geq 0 \quad (2).$$

Für dieses x würde dann folgen, dass

$$0 \stackrel{(1)}{\geq} 2x + 1 = 2x - 1 + 2 \stackrel{(2)}{\geq} 0 + 2 = 2.$$

Dies ist aber ein Widerspruch. Also war unsere Annahme falsch und somit die zu beweisende Aussage richtig. $\square$

Das dabei verwendete Symbol „ $\square$ “ ist die übliche Art, das Ende eines Beweises zu kennzeichnen. Zur Verdeutlichung haben wir die beiden Ungleichungen mit (1) und (2) markiert, um später angeben zu können, wo sie verwendet werden.

- (b) Manchmal weiß man von einer Aussage aufgrund der Aufgabenstellung zunächst einmal noch nicht, ob sie wahr oder falsch ist. In diesem Fall muss man sich dies natürlich zuerst überlegen – und, falls die Aussage falsch ist, ihre Negation beweisen. Als Beispiel dafür betrachten wir die Aufgabe

Man beweise oder widerlege: Für alle $x \in \mathbb{R}$ gilt $2x + 1 < 0$ oder $2x - 1 > 0$.

In diesem Fall merkt man schnell, dass die Aussage falsch sein muss, weil die Ungleichungen schon für den Fall $x = 0$ nicht stimmen. Man könnte als Lösung der Aufgabe unter Beachtung der Negationsregeln aus Bemerkung 1.8 also aufschreiben:

Behauptung: Die Aussage ist falsch, d. h. es gibt ein $x \in \mathbb{R}$ mit $2x + 1 \geq 0$ und $2x - 1 \leq 0$.

Beweis: Für $x = 0$ ist $2x + 1 = 1 \geq 0$ und $2x - 1 = -1 \leq 0$. $\square$

Beachte, dass dies ein vollständiger Beweis ist: *Um eine allgemeine Aussage zu widerlegen, genügt es, ein Gegenbeispiel dafür anzugeben.*

Bemerkung 1.11. Bevor wir unsere kurze Auflistung der für uns wichtigen Prinzipien der Logik beenden, wollen wir noch kurz auf ein paar generelle Dinge eingehen, die man beim Aufschreiben mathematischer Beweise oder Rechnungen beachten muss.

Dass wir bei unseren logischen Argumenten sauber und exakt arbeiten – also z. B. nicht Folgerungen, die keine Äquivalenzen sind, in der falschen Richtung verwenden, „für alle“ mit „es gibt“ verwechseln oder Ähnliches – sollte sich von selbst verstehen. Die folgende kleine Geschichte hilft vielleicht zu verstehen, was damit gemeint ist.

Ein Ingenieur, ein Physiker und ein Mathematiker fahren mit dem Zug nach Frankreich und sehen dort aus dem Fenster des Zuges ein schwarzes Schaf.

Da sagt der Ingenieur: „Oh, in Frankreich sind die Schafe schwarz!“

Darauf der Physiker: „Nein . . . wir wissen jetzt nur, dass es in Frankreich mindestens ein schwarzes Schaf gibt.“

Der Mathematiker: „Nein . . . wir wissen nur, dass es in Frankreich mindestens ein Schaf gibt, das auf mindestens einer Seite schwarz ist.“

Es gibt aber noch einen weiteren sehr wichtigen Punkt, der leider oft nicht beachtet wird: In der Regel werden wir beim Aufschreiben sowohl Aussagen notieren wollen, die wir erst noch zeigen wollen (um schon einmal zu sagen, worauf wir hinaus wollen), als auch solche, von denen wir bereits wissen, dass sie wahr sind (z. B. weil sie für die zu zeigende Behauptung als wahr vorausgesetzt werden oder weil sie sich logisch aus irgendetwas bereits Bekanntem ergeben haben). Es sollte offensichtlich sein, dass wir Aussagen mit derartig verschiedenen Bedeutungen für die Argumentationsstruktur nicht einfach zusammenhangslos hintereinander schreiben dürfen, wenn noch jemand in der Lage sein soll, die Argumente nachzuvollziehen. Betrachten wir z. B. noch einmal unseren

Beweis aus Beispiel 1.10 (a) oben, so wäre eine Art des Aufschreibens in folgendem Stil (wie man es leider oft sieht)

$$\begin{aligned} 2x + 1 &> 0 \text{ oder } 2x - 1 < 0 \\ 2x + 1 &\leq 0 \quad \quad 2x - 1 \geq 0 \\ 0 &\geq 2x + 1 = 2x - 1 + 2 \geq 0 + 2 = 2 \end{aligned}$$

völlig inakzeptabel, obwohl hier natürlich letztlich die gleichen Aussagen stehen wie oben. Kurz gesagt:

Von *jeder* aufgeschriebenen Aussage muss für den Leser *sofort* und *ohne eigenes Nachdenken* ersichtlich sein, welche Rolle sie in der Argumentationsstruktur spielt: Ist es z. B. eine noch zu zeigende Behauptung, eine Annahme oder eine Folgerung (und wenn ja, aus was)?

Dies bedeutet allerdings nicht, dass wir ganze Aufsätze schreiben müssen. Eine (schon recht platz-optimierte) Art, den Beweis aus Beispiel 1.10 (a) aufzuschreiben, wäre z. B.

Angenommen, es gäbe ein $x \in \mathbb{R}$ mit $2x + 1 \leq 0$ und $2x - 1 \geq 0$.

Dann wäre $0 \geq 2x + 1 = 2x - 1 + 2 \geq 0 + 2 = 2$, Widerspruch. □

01

1.B Mengenlehre

Nachdem wir die wichtigsten Regeln der Logik behandelt haben, wenden wir uns jetzt der Mengenlehre zu. Die gesamte moderne Mathematik basiert auf diesem Begriff der Menge, der ja auch schon aus der Schule hinlänglich bekannt ist. Zur Beschreibung, was eine Menge ist, zitiert man üblicherweise die folgende Charakterisierung von Georg Cantor (1845–1918):

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

Die in einer Menge zusammengefassten Objekte bezeichnet man als ihre **Elemente**.

Notation 1.12.

- (a) Wir schreiben $x \in M$, falls x ein Element der Menge M ist, und $x \notin M$ andernfalls.
- (b) Die einfachste Art, eine Menge konkret anzugeben, besteht darin, ihre Elemente in geschweiften Klammern aufzulisten, wobei es auf die Reihenfolge und Mehrfachnennungen nicht ankommt. So sind z. B. $\{1, 2, 3\}$ und $\{2, 3, 1, 3\}$ zwei Schreibweisen für dieselbe Menge mit den drei Elementen 1, 2 und 3.

Beachte, dass die Elemente einer Menge nicht unbedingt Zahlen sein müssen – so ist z. B. $M = \{\{2, 3\}, \{1, 3\}\}$ eine Menge mit zwei Elementen, die selbst wieder Mengen sind, nämlich $\{2, 3\}$ und $\{1, 3\}$. Mit der Notation aus (a) ist also z. B. $\{1, 3\} \in M$. Insbesondere ist M nicht dasselbe wie die Menge $\{1, 2, 3\}$.

- (c) Man kann die Elemente einer Menge auch durch eine beschreibende Eigenschaft angeben: $\{x : A(x)\}$ bezeichnet die Menge aller Objekte x , für die die Aussage $A(x)$ wahr ist, wie z. B. in $\{x \in \mathbb{R} : x^2 = 1\} = \{-1, 1\}$.
- (d) Die Menge $\{\}$ ohne Elemente, die sogenannte **leere Menge**, bezeichnen wir mit $\emptyset$.
- (e) Eine Menge M heißt **Teilmenge** einer Menge N (geschrieben $M \subset N$), wenn jedes Element von M auch Element von N ist, bzw. in der Quantorenschreibweise von Notation 1.6 wenn

$$\forall x: x \in M \Rightarrow x \in N.$$

Man sagt in diesem Fall auch, dass N eine **Obermenge** von M ist (geschrieben $N \supset M$).

Beachte, dass M und N dabei auch gleich sein können; in der Tat ist offensichtlich

$$M = N \quad \text{genau dann, wenn} \quad M \subset N \text{ und } N \subset M.$$

Oft wird man eine Gleichheit $M = N$ von Mengen auch so beweisen, dass man separat $M \subset N$ und $N \subset M$ zeigt.

Wenn wir ausdrücken wollen, dass M eine Teilmenge von N und nicht gleich N ist, so schreiben wir dies als $M \subsetneq N$ und sagen, dass M eine **echte Teilmenge** von N ist. Es ist wichtig, dies von der Aussage $M \not\subset N$ zu unterscheiden, die bedeutet, dass M keine Teilmenge von N ist.

Achtung: Manchmal wird in der Literatur das Symbol „ $\subset$ “ für *echte* Teilmengen und „ $\subseteq$ “ für nicht notwendig echte Teilmengen verwendet.

- (f) Hat eine Menge M nur endlich viele Elemente, so nennt man M eine **endliche Menge** und schreibt die Anzahl ihrer Elemente als $|M|$. Andernfalls setzt man formal $|M| = \infty$.

Bemerkung 1.13 (Russellsches Paradoxon). Die oben gegebene Charakterisierung von Mengen von Cantor ist aus mathematischer Sicht natürlich sehr schwammig. In der Tat hat Bertrand Russell kurz darauf bemerkt, dass sie sogar schnell zu Widersprüchen führt. Er betrachtet dazu

$$M = \{A : A \text{ ist eine Menge mit } A \notin A\}, \quad (*)$$

also „die Menge aller Mengen, die sich nicht selbst als Element enthalten“. Sicherlich ist es eine merkwürdige Vorstellung, dass eine Menge sich selbst als Element enthalten könnte – im Sinne von Cantors Charakterisierung wäre die Definition (*) aber zulässig. Fragen wir uns nun allerdings, ob sich die so konstruierte Menge M selbst als Element enthält, so erhalten wir sofort einen Widerspruch: Wenn $M \in M$ gilt, so würde das nach der Definition (*) ja gerade bedeuten, dass $M \notin M$ ist – und das wiederum, dass doch $M \in M$ ist. Man bezeichnet dies als das *Russellsche Paradoxon*.

Die Ursache für diesen Widerspruch ist, dass die Definition (*) rückbezüglich ist: Wir wollen eine neue Menge M konstruieren, verwenden dabei aber auf der rechten Seite der Definition *alle Mengen*, also u. a. auch die Menge M , die wir gerade erst definieren wollen. Das ist in etwa so, als würdet ihr im Beweis eines Satzes die Aussage des Satzes selbst verwenden – und das ist natürlich nicht zulässig.

Man muss bei der Festlegung, was Mengen sind und wie man sie bilden kann, also eigentlich viel genauer vorgehen, als es Cantor getan hat. Heutzutage verwendet man hierzu in der Regel das im Jahre 1930 aufgestellte Axiomensystem von Zermelo und Fraenkel, das genau angibt, wie man aus bekannten Mengen neue konstruieren darf: z. B. indem man sie schneidet oder vereinigt, oder aus bereits bekannten Mengen Elemente mit einer bestimmten Eigenschaft auswählt. Wir wollen dies hier in dieser Vorlesung aber nicht weiter thematisieren und uns mit der naiven Mengencharakterisierung von Cantor begnügen (sowie der Versicherung meinerseits, dass schon alles in Ordnung ist, wenn wir neue Mengen immer nur aus alten konstruieren und keine rückbezüglichen Definitionen hinschreiben). Genauer zum Zermelo-Fraenkel-Axiomensystem könnt ihr z. B. in [E, Kapitel 13] nachlesen.

Notation 1.14 (Reelle Zahlen). Unser wichtigstes Beispiel für eine Menge ist die Menge der **reellen Zahlen**, die wir mit $\mathbb{R}$ bezeichnen werden. Wir wollen die Existenz der reellen Zahlen in dieser Vorlesung axiomatisch voraussetzen und begnügen uns daher an dieser Stelle damit zu sagen, dass man sie sich als die Menge der Punkte auf einer Geraden (der „Zahlengeraden“) vorstellen kann. Zusätzlich werden wir in den nächsten beiden Kapiteln die mathematischen Eigenschaften von $\mathbb{R}$ exakt angeben (und ebenfalls axiomatisch voraussetzen) – und zwar genügend viele Eigenschaften, um $\mathbb{R}$ dadurch eindeutig zu charakterisieren.

Ich möchte hier noch einmal betonen, dass man die Existenz und die Eigenschaften der reellen Zahlen eigentlich nicht voraussetzen müsste: Man kann das auch allein aus den Axiomen der Logik und Mengenlehre herleiten! Dies wäre jedoch relativ aufwendig und würde euch im Moment mehr verwirren als helfen, daher wollen wir hier darauf verzichten. Wer sich trotzdem dafür interessiert, kann die Einzelheiten hierzu in [E, Kapitel 1 und 2] nachlesen.

Außer den reellen Zahlen sind vor allem noch die folgenden Teilmengen von $\mathbb{R}$ wichtig:

- (a) die Menge $\mathbb{N} = \{0, 1, 2, \dots\}$ der **natürlichen Zahlen** (Achtung: In der Literatur wird die 0 manchmal nicht mit zu den natürlichen Zahlen gezählt!);

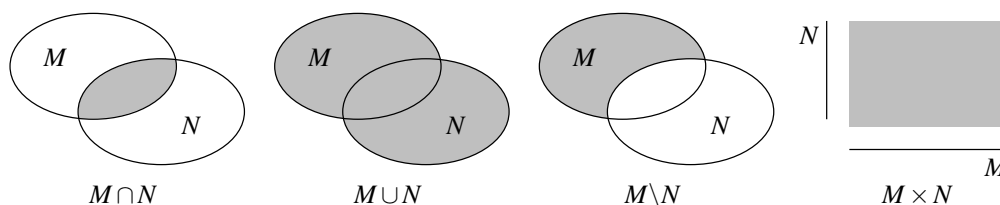
- (b) die Menge $\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}$ der **ganzen Zahlen**;
 (c) die Menge $\mathbb{Q} = \{\frac{p}{q} : p, q \in \mathbb{Z}, q \neq 0\}$ der **rationalen Zahlen**.

Offensichtlich sind diese Mengen ineinander enthalten: Es gilt $\mathbb{N} \subsetneq \mathbb{Z} \subsetneq \mathbb{Q} \subsetneq \mathbb{R}$. Teilmengen von $\mathbb{R}$, die durch Ungleichungen gegeben sind, schreiben wir in der Regel, indem wir die Ungleichungsbedingung als Index an das Symbol $\mathbb{R}$ schreiben, z. B. $\mathbb{R}_{\geq 0}$ für die Menge $\{x \in \mathbb{R} : x \geq 0\}$ aller nicht-negativen Zahlen.

Notation 1.15. Sind M und N Mengen, so bezeichnen wir mit ...

- (a) $M \cap N := \{x : x \in M \text{ und } x \in N\}$ die **Schnittmenge** von M und N . Gilt $M \cap N = \emptyset$, so sagen wir, dass M und N **disjunkt** sind.
 (b) $M \cup N := \{x : x \in M \text{ oder } x \in N\}$ die **Vereinigungsmenge** von M und N . Im Fall einer **disjunkten Vereinigung** mit $M \cap N = \emptyset$ schreiben wir statt $M \cup N$ auch $M \sqcup N$.
 (c) $M \setminus N := \{x : x \in M \text{ und } x \notin N\}$ die **Differenzmenge** von M und N .
 (d) $M \times N := \{(x, y) : x \in M, y \in N\}$ die **Produktmenge** bzw. das Produkt von M und N . Die Schreibweise (x, y) steht hierbei für ein **geordnetes Paar**, d. h. einfach für die Angabe eines Elements aus M und eines aus N (wobei es auch im Fall $M = N$ auf die Reihenfolge ankommt, d. h. (x, y) ist genau dann gleich (x', y') wenn $x = x'$ und $y = y'$). Im Fall $M = N$ schreibt man $M \times N = M \times M$ auch als M^2 .
 (e) $\mathcal{P}(M) := \{A : A \text{ ist Teilmenge von } M\}$ die **Potenzmenge** von M .

Das Symbol „ $:=$ “ bedeutet hierbei, dass der Ausdruck auf der linken Seite durch die rechte Seite definiert wird. Die Konstruktionen (a) bis (d) können durch die folgenden Bilder veranschaulicht werden. Natürlich sind sie auch für mehr als zwei Mengen möglich; aus der Schule kennt ihr zum Beispiel sicher den Fall $\mathbb{R}^3 = \mathbb{R} \times \mathbb{R} \times \mathbb{R}$.



Die Potenzmenge $\mathcal{P}(M)$ aller Teilmengen einer gegebenen Menge M lässt sich dagegen nicht so einfach durch ein Bild darstellen. Es ist z. B.

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

Aufgabe 1.16. Wie lautet die Negation der folgenden Aussagen? Formuliere außerdem die Aussage (a) in Worten (also analog zu (b)) sowie die Aussage (b) mit Quantoren und anderen mathematischen Symbolen (also analog zu (a)).

- (a) $\forall n \in \mathbb{N} \exists m \in \mathbb{N} : n = 2m$.
 (b) Zwischen je zwei verschiedenen reellen Zahlen gibt es noch eine weitere reelle Zahl.
 (c) Sind M, N, R Mengen mit $R \subset N \subset M$, so ist $M \setminus N \subset M \setminus R$.

Aufgabe 1.17. Es seien A, B, C Aussagen und M, N, R Mengen. Man zeige:

- (a) $A \vee (B \wedge C) \Leftrightarrow (A \vee B) \wedge (A \vee C)$ und $A \wedge (B \vee C) \Leftrightarrow (A \wedge B) \vee (A \wedge C)$.
 (b) $M \cup (N \cap R) = (M \cup N) \cap (M \cup R)$ und $M \cap (N \cup R) = (M \cap N) \cup (M \cap R)$.

Aufgabe 1.18. Welche der folgenden Aussagen sind für beliebige gegebene x und M äquivalent zueinander? Zeige jeweils die Äquivalenz bzw. widerlege sie durch ein Gegenbeispiel.

- (a) $x \in M$ (b) $\{x\} \subset M$ (c) $\{x\} \cap M \neq \emptyset$
 (d) $\{x\} \in M$ (e) $\{x\} \setminus M = \emptyset$ (f) $M \setminus \{x\} = \emptyset$

Aufgabe 1.19. Man beweise oder widerlege: Für alle Mengen $A \subset M$ und $A' \subset M'$ gibt es Teilmengen B und C von M sowie B' und C' von M' , so dass

$$(M \times M') \setminus (A \times A') = (B \times B') \cup (C \times C').$$

Könnt ihr die Aussage durch eine Skizze veranschaulichen?