

Kapitel 1
Grundlagen

1

1

1	Grundlagen	
1.1	Mengenbegriff	7
1.1.1	Relationen zwischen Mengen	8
1.1.2	Das kartesische Produkt	9
1.2	Elemente der Logik	9
1.2.1	Aussagen	10
1.2.2	Quantoren	12
1.2.3	Logische Argumente	13
1.3	Zahlssysteme und elementares Rechnen	18
1.3.1	Die natürlichen Zahlen	18
1.3.2	Die ganzen Zahlen	19
1.3.3	Die rationalen Zahlen (Bruchzahlen)	19
1.3.4	Die reellen Zahlen	20
1.4	Potenzen, Wurzeln	22
1.4.1	Motivation	22
1.4.2	Potenzen	22
1.4.3	Wurzeln	23
1.4.4	Lösen von Potenzgleichungen	23
1.4.5	Prozentrechnung, Rechnen mit Wachstumsraten	25
1.5	Kombinatorik	27
1.6	Reelle Zahlenfolgen	33
1.6.1	Motivation	33
1.6.2	Begriffsbildung	33
1.7	Reihen	36
1.7.1	Motivation	36
1.7.2	Summen (Endliche Reihen)	38
1.7.3	Unendliche Reihen	38
1.7.4	Die (endliche) geometrische Reihe	39
1.8	Funktionen und Abbildungen	40
1.8.1	Komposition von Funktionen	44
1.8.2	Umkehrfunktion	45
1.9	Stetigkeit	46
1.9.1	Motivation	46
1.9.2	Begriffsbildung	47
1.9.3	Eigenschaften stetiger Funktionen	48
1.10	Exponentialfunktion	49
1.10.1	Definition	49
1.10.2	Eigenschaften	50
1.11	Kontinuierliches Wachstum	51
1.12	Der Logarithmus	52
1.12.1	Rechenregeln	53

1 Grundlagen

1.1 Mengenbegriff

1.1

Eine **Menge** ist eine Zusammenfassung bestimmter, wohlunterschiedener **Objekte** unserer Anschauung oder unseres Denkens zu einem Ganzen. Die Objekte der Menge heißen **Elemente** der Menge. Aus dieser Definition ergibt sich, dass eine Menge bestimmt ist, wenn man alle Elemente angegeben hat, die zur Menge gehören.

Beispiele für Mengen:

1. Die Menge aller Säugetiere.
2. Die Menge aller Einsetzungen x , so dass „ x ist ein Einzeller“ eine wahre Aussage ist: $\{x : x \text{ ist Einzeller}\}$.
3. $\mathbb{N} = \{1, 2, 3, \dots\}$ ist die Menge der natürlichen Zahlen.
4. $P = \{x \in \mathbb{N} : x \text{ ist eine Primzahl}\}$.
5. $L = \{x \in \mathbb{N} : x^2 = 4\} = \{2\}$.

Die Beispiele illustrieren, dass man Mengen auf verschiedene Weisen angeben kann: Durch Aufzählung der Elemente, z.B.

$$A = \{1, 2, 3, 4\},$$

durch Angabe einer sie charakterisierenden Eigenschaft,

$$A = \{x \in \mathbb{N} : 1 \leq x \leq 4\},$$

bzw.

„ A ist die Menge aller natürlichen Zahlen von 1 bis einschließlich 4“,

oder durch eine graphische Darstellung. Man verwendet bei der Angabe von Mengen die geschweiften Mengenklammern $\{$ und $\}$. Bei einer Aufzählung wie $\{a, b, c\}$ kommt es nicht auf die Reihenfolge der Elemente an. Das heißt, die Mengen $\{a, b, c\}$ und $\{c, a, b\}$ bezeichnen dieselben Mengen.

Mengen können beliebige Elemente enthalten, auch wieder Mengen. $A = \{\{1, 2\}, \{3, 4\}\}$ ist die Menge, welche die zwei Elemente $\{1, 2\}$ und $\{3, 4\}$ enthält.

Die **leere Menge**, die kein Element enthält, wird mit $\emptyset$ oder auch $\{\}$ bezeichnet.

Man schreibt $a \in A$, wenn a Element der Menge A ist. Ansonsten schreibt man $a \notin A$.

Beispiele: 1 ist Element der Menge $\{1, 3, 5\}$.

► 1.1.1 Relationen zwischen Mengen

▷ Teilmengen

Sind A, B zwei Mengen, so heißt A **Teilmenge** von B , i.Z. $A \subseteq B$, wenn jedes Element $a \in A$ auch in B enthalten ist, also wenn gilt: Aus $a \in A$ folgt $a \in B$. Gibt es Elemente in B , die nicht in A enthalten sind, so ist A eine echte Teilmenge und man schreibt: $A \subset B$.

Beispiele: (i) $A = \{1, 2\}$ ist eine Teilmenge von $\mathbb{N}$: $A \subset \mathbb{N}$. (ii) Da alle Katzen Säugetiere sind, ist die Menge aller Katzen eine Teilmenge der Menge aller Säugetiere.

▷ Gleichheit von Mengen

Zwei Mengen A und B sind gleich, i.Z. $A = B$, wenn sie dieselben Elemente enthalten. $A = B$ ist also gleichbedeutend mit der simultanen Gültigkeit von

$$A \subseteq B : \text{Wenn } x \in A, \text{ dann auch } x \in B$$

und

$$B \subseteq A : \text{Wenn } x \in B, \text{ dann auch } x \in A.$$

Um die Gleichheit von zwei Mengen zu verifizieren, zeigt man, dass jedes Element der Menge A auch in B enthalten ist, und umgekehrt.

So sind die Mengen $A = \{2, 7^2\}$ und $B = \{\sqrt{4}, 49\}$ gleich: Es ist $2 = \sqrt{4} \in B$ und $7^2 = 49 \in B$. Umgekehrt ist $\sqrt{4} = 2 \in A$ und $49 = 7 \cdot 7 = 7^2 \in A$.

▷ Durchschnitt

Der **Durchschnitt** (Schnitt) von zwei Mengen A und B ist gegeben durch

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

A und B heißen **disjunkt**, wenn ihr Schnitt leer ist, d.h. wenn $A \cap B = \emptyset$.

▷ Vereinigung

Die **Vereinigungsmenge** von A und B ist gegeben durch

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

Man spricht von einer *disjunkten* Vereinigung, wenn $A \cap B = \emptyset$.

▷ Distributivgesetze

Es gelten die Distributivgesetze der Durchschnitts- und Vereinigungsbildung:

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

sowie

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

Man mache sich die Gültigkeit dieser Regeln an Venn-Diagrammen (Mengenkreisen) klar!

➤ 1.1.2 Das kartesische Produkt

Sind A und B zwei Mengen, so ist das **kartesische Produkt** die Menge aller **2-Tupel** (a, b) mit $a \in A$ (1. Koordinate oder Komponente) und $b \in B$ (2. Koordinate). Formal:

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

Im Gegensatz zu Mengen kommt es bei Tupeln auf die *Reihenfolge* der Elemente an. Zwei Mengen sind gleich, wenn ihre Elemente übereinstimmen. Zwei 2-Tupel sind gleich, wenn sowohl die 1. Koordinate als auch 2. Koordinate übereinstimmen.

Beispiel 1.1.1 Zeichnet man eine Gerade G mit Steigung b und y -Achsenabschnitt a in ein Koordinatensystem ein, so kann man die Gerade durch die Menge aller Punkte (x, y) beschreiben, die auf der Geraden liegen: **1.1.1**

$$G = \{(x, y) \in \mathbb{R} \times \mathbb{R} : y = a + b \cdot x\}.$$

Hierbei bezeichnet $\mathbb{R}$ die Menge aller reellen Zahlen, dazu später mehr.

Beispiel 1.1.2 Man plant, ein Experiment dreimal zu wiederholen. Jedes einzelne Experiment bestehe darin, die elektrische Leitfähigkeit (in Ohm) eines Blattes zu messen. Das Experiment kann nun durch die Menge aller möglichen Messwert-Paare beschrieben werden, die man prinzipiell erhalten kann: **1.1.2**

$$E = \{(x_1, x_2, x_3) : 0 \leq x_1 < \infty, 0 \leq x_2 < \infty, 0 \leq x_3 < \infty\}.$$

Bei dem Tripel (x_1, x_2, x_3) steht x_i für das Ergebnis des i -ten Experiments, $i = 1, 2, 3$.

1.2 Elemente der Logik

1.2

Die Logik beschäftigt sich mit dem Wahrheitsgehalt von Aussagen und der Korrektheit und Schlüssigkeit von Argumenten.

► 1.2.1 Aussagen

Im Sinne der Logik versteht man unter einer **Aussage** ein sprachliches Gebilde, das entweder wahr (W) oder falsch (F) ist. Da wir hier lediglich einige wesentliche Grundlagen betrachten wollen, beschränken wir uns auf Aussagen, bei denen der Wahrheitswert (W oder F) nicht vom Kontext abhängt. Somit sind die Sätze wie

1. Alle Katzen sind Raubtiere.
2. Mitochondrien-DNA vererbt sich ohne Rekombination.
3. Bochum liegt im Ruhrgebiet.
4. Wasserstoff und Sauerstoff reagieren zu Wasser.

Aussagen im Sinne der Logik. Sätze wie

1. Ich finde das heutige Fernsehprogramm langweilig.
2. Guten Abend!
3. Man kann nicht durch das Brandenburger Tor gehen.

sind jedoch keine Aussagen in unserem (vereinfachenden) Sinne. Der erste Satz ist eine reine Meinungsäußerung. Der „Wahrheitswert“ hängt von der Person ab, die ihn sagt. Der zweite Satz ist ein Ausruf, dem man keinen Wahrheitswert zuordnen kann. Beim letzten Satz hängt der Wahrheitswert vom Zeitpunkt, also vom Kontext ab. 1976 war der Satz wahr, 2001 jedoch falsch. Natürlich gibt es Grenzfälle, doch der Einfachheit halber sprechen wir hier nur dann von einer Aussage, wenn wir den Wahrheitswert eindeutig feststellen können, wenn wir die Bedeutung (Semantik) der einzelnen Begriffe kennen.

Aussagen sollen im Folgenden mit Großbuchstaben $A, B, \dots$ abgekürzt werden.

► Verknüpfen von Aussagen durch Junktoren

Zu den universellen Bestandteilen von Aussagen gehören Worte (Wortkomplexe) wie „und“, „oder“, „entweder...oder...“, „nicht“, „wenn-dann“, mit deren Hilfe Aussagen verknüpft werden. Solche Worte heißen Junktoren. So werden die Aussagen „ $3 < 4$ “ und „ $3 = 4$ “ durch den ODER-Junktor zu „ $3 \leq 4$ “ verknüpft.

Junktoren gibt man an, indem man ihren Wahrheitswert für alle möglichen Belegungen der Aussagen mit Wahrheitswerten in einer Wahrheitstafel aufschreibt. Man spricht von einem n -stelligen Junktor, wenn der Junktor n Aussagen (die Argumente) verknüpft. Ein n -stelliger Junktor entspricht also einer Wahrheitstafel mit 2^n Einträgen.

Negation ($\neg$): Die Negation ist ein 1-stelliger Junktoren, der den Wahrheitswert einer Aussage umkehrt.

A	$\neg A$
W	F
F	W

Konjunktion UND ($\wedge$): $A \wedge B$ ist wahr, wenn sowohl A als auch B wahr sind. Sonst ist $A \wedge B$ falsch.

A	B	$A \wedge B$
W	W	W
W	F	F
F	W	F
F	F	F

Disjunktion ODER ($\vee$): Das Wort „oder“ wird in der Umgangssprache in zweierlei Weise gebraucht. Das sog. inklusive (einschließende) ODER ist auch dann wahr, wenn beide Aussagen wahr sind, während das exklusive (ausschließende) ODER (entweder-oder) diesen Fall mit dem Wahrheitswert F belegt. Wir vereinbaren, dass $A \vee B$ das inklusive ODER darstellen soll, also falsch ist, wenn sowohl A als auch B falsch sind, und in jedem anderen Fall wahr ist. Für das exklusive ODER verwenden wir das Symbol XOR.

A	B	$A \vee B$	$A \text{ XOR } B$
W	W	W	F
W	F	W	W
F	W	W	W
F	F	F	F

Implikation WENN-DANN ($\Rightarrow$): Die Folgerungsbedingung $A \Rightarrow B$ ist falsch, wenn A wahr, B jedoch falsch ist. In allen anderen Fällen ist $A \Rightarrow B$ wahr.

A	B	$A \Rightarrow B$	$B \Rightarrow A$	$A \Leftrightarrow B$
W	W	W	W	W
W	F	F	W	F
F	W	W	F	F
F	F	W	W	W

Bei einer Implikation $A \Rightarrow B$ sagt man, B sei **notwendige Bedingung** für A , da B notwendigerweise wahr ist, wenn A wahr ist. A heißt **hinreichende Bedingung** für B , da die Wahrheit von A ausreicht, um die Wahrheit von B zu erzwingen.

Äquivalenz ($\Leftrightarrow$): $A \Leftrightarrow B$ ist genau dann wahr, wenn sowohl $A \Rightarrow B$ als auch $B \Rightarrow A$ wahr sind.

$A \Leftrightarrow B$ ist also genau dann wahr, wenn A und B denselben Wahrheitswert haben: A und B sind *gleichbedeutend*.

▷ Wichtige Äquivalenzen

(1) $A \Rightarrow B$ ist äquivalent zu $\neg B \Rightarrow \neg A$.

„Wenn Sauerstoff (O_2) und Wasserstoff ($2H_2$) zusammenkommen, dann entsteht Wasser ($2H_2O$). Das heißt: Ist kein Wasser entstanden, dann sind H_2 und O_2 auch nicht zusammengekommen.“

„Wenn man die Bremse betätigt, so hält das Auto an. Das heißt: Wenn das Auto nicht anhält, so ist die Bremse nicht betätigt.“

(2) $\neg(A \wedge B)$ ist äquivalent zu $(\neg A) \vee (\neg B)$.

(3) $A \Rightarrow B$ ist äquivalent zu $(\neg A) \vee B$.

► 1.2.2 Quantoren

Betrachten wir die folgende wahre Aussage:

A : Bei Mäusen wird der Nachwuchs von den Weibchen gestillt.

Diese Aussage kann offensichtlich verallgemeinert werden. Sie ist nicht nur für Mäuse gültig, sondern überhaupt für alle Säugetiere. Nun ist es natürlich unschön, diese Aussage für alle Säugetierarten formulieren zu müssen. Viel eleganter ist es, einen *Platzhalter* einzuführen, der den Begriff *Mäuse* ersetzt. Einen solchen Platzhalter nennt man **Variable**. Variablen werden in der Regel durch die Buchstaben x, y oder z bezeichnet (dies ist kein Gesetz, sondern eine Konvention). Hierdurch erhalten wir eine **Aussageform**:

$A(x)$: Bei x wird der Nachwuchs von den Weibchen gestillt.

Je nachdem, was man für x einsetzt, erhält man eine wahre, falsche oder sinnlose Aussage. Bezeichnen wir mit S die Menge aller Säugetierarten, so können wir formulieren:

Für alle $x \in S$ gilt: Bei x wird der Nachwuchs von den Weibchen gestillt.

Der **Allquantor** $\forall$ verallgemeinert eine Aussageform auf alle Einsetzungen einer Menge:

$$\forall x \in G : A(x).$$

Dies ist so zu lesen: Für alle $x \in G$ gilt $A(x)$. Man erhält immer dann eine wahre Aussage, wenn man ein Element aus der Menge G einsetzt. Man nennt dies eine *Allaussage* (Generalisierung). Die Allaussage ist z.B. richtig, wenn G die Menge aller Mäuse ist.

Ein mathematisches Beispiel: Für alle $0 \leq x \leq 2$ gilt: $x^2 \leq 4$. Die Aussageform $A(x) : x^2 \leq 4$ ist also immer dann richtig, wenn man ein x einsetzt, welches die Bedingung $0 \leq x \leq 2$ erfüllt, Kurzform:

$$\forall 0 \leq x \leq 2 : x^2 \leq 4$$

oder auch: $\forall x \in [0, 2] : x^2 \leq 4$.

Der Existenzquantor $\exists$ postuliert die Wahrheit von $A(x)$ für zumindest ein x . Bei einer solchen Existenzaussage (Partikularisierung) wird also die Existenz einer Einsetzungsmöglichkeit für die Variable x behauptet, so dass man eine wahre Aussage erhält. Betrachten wir die Aussageform

$A(x)$: Es gibt ein Tier $x \in T$, das fliegen kann.

Für Einsetzungen aus der Menge T aller Vögel ist das offenkundig richtig, nimmt man für T die Menge aller Fische, dann bleibt es richtig (Beispiel: *exocoetus volitans*), ist T hingegen die Menge aller Löwen, so erhält man nie eine wahre Aussage.

Erwähnenswert sind noch die folgenden Verneinungsregeln:

$$(1) \neg(\forall x \in G : A(x)) \Leftrightarrow \exists x \in G : \neg A(x)$$

Die Aussage „Für alle $x \in G$ gilt: $A(x)$ “ ist genau dann falsch, wenn es ein Gegenbeispiel $x \in G$ gibt, für das $A(x)$ gilt.

$$(2) \neg(\exists x \in G : A(x)) \Leftrightarrow \forall x \in G : \neg A(x)$$

Die Aussage „Es gibt ein $x \in G$ mit $A(x)$ “ ist genau dann falsch, wenn für alle $x \in G$ gilt: $A(x)$ gilt nicht.

► 1.2.3 Logische Argumente

Im Sinne der Logik sind Argumente Folgen (Aneinanderreihungen) von Aussagen, die das Ziel verfolgen, eine Folgerungsbeziehung zwischen den **Prämissen** (Annahmen, Bedingungen) und der **Konklusion** (Folgerung, Schlussfolgerung) rational (logisch) erscheinen zu lassen. Dies erfolgt durch eine lückenlose Rückführung auf bereits anerkannte Aussagen. Ist das Argument korrekt, so ist die Folgerungsbeziehung logisch zwingend im Sinne der Implikation $\Rightarrow$.

„Da der Kampf gegen Nachbarn ein Übel ist und der Kampf gegen die Thebaner ein Kampf gegen Nachbarn ist, ist es klar, dass der Kampf gegen die Thebaner ein Übel ist.“ (Sokrates)

Ein Argument wird im Sinne der Logik in seine **Normalform** überführt, indem man die Prämissen explizit untereinander schreibt und die Konklusion durch ein „Also:“ kenntlich macht. Hierzu muss man u.U. die Prämissen aus dem Text rekonstruieren. Die Normalform des obigen Beispiels lautet also:

1. Der Kampf gegen Nachbarn ist ein Übel.
 2. Der Kampf gegen die Thebaner ist ein Kampf gegen Nachbarn.
- Also: Der Kampf gegen die Thebaner ist ein Übel.

Betrachten wir noch ein historisches Beispiel:

„Die natürliche Zuchtwahl wählt die Besten aus. Wäre das nicht der Fall, könnte die Erde innerhalb weniger Jahrhunderte nicht mehr die Nachkommenschaft eines einzigen Paares fassen.“ (Charles Darwin)

Um die Normalform aufzustellen, muss man einige Prämissen ergänzen, um zu einer klaren logischen Schlusskette zu gelangen. Dies sollte mit einiger Umsicht entlang des Originaltextes erfolgen.

1. Das evolutionäre Ziel ist die Erhaltung der Art.
 2. Der Lebensraum auf der Erde ist begrenzt.
 3. Es gibt einen Geburtenüberschuss.
 4. Gibt es einen Geburtenüberschuss und keine natürliche Auswahl der Besten, so wächst eine Art über alle Grenzen.
 5. Wenn der Lebensraum begrenzt ist und eine Art über alle Grenzen wächst, dann zerstört sie ihr Existenzgrundlage.
- Also: Durch die natürliche Zuchtwahl werden die Besten ausgewählt.

Anmerkung: Heute geht man nicht mehr davon aus, dass Arterhaltung ein evolutionäres Ziel ist (Stichworte: Soziobiologie, egoistisches Gen). „Natürliche Zuchtwahl“ und „Auswahl der Besten“ sind bekanntlich auch falsche Übersetzungen. Wie sollte Darwins Argument biologisch korrekt formuliert werden?

▷ Korrektheit und Schlüssigkeit von Argumenten

Ein Argument ist **(formal) korrekt**, wenn die Konklusion wahr ist, immer dann wenn alle Prämissen wahr sind. Genauer muss man sagen: immer dann, wenn alle Prämissen als wahr angenommen werden. Ist eine der Prämissen falsch, so ist der Schluss (trotzdem) formal korrekt.

Im folgenden Beispiel sind zwar Konklusion und Prämisse wahr, aber die Prämisse ist keine Begründung für die Konklusion.

Im Jahr 79 wurde Pompeji durch den Ausbruch des Vesuv zerstört.

Also: Albert Einstein starb im Jahr 1955 in Princeton.

Korrektheit ist eine notwendige Bedingung für ein gutes Argument. Formale Korrektheit bedeutet aber nur, dass die Konklusion wahr ist, wenn man die Prämissen als wahr annimmt. Das folgende Argument ist zwar formal korrekt, aber sowohl die zweite Prämisse als auch die Konklusion sind in offenkundiger

Weise falsch, wenn man den üblichen Sprachgebrauch der Begriffe zugrunde legt.

1. Alle Wale sind Fische.
 2. Alle Delphine sind Wale.
- Also: Alle Delphine sind Fische.

Ein Argument heißt **schlüssig**, wenn es korrekt ist und wenn alle seine Prämissen wahr sind. Beispiel:

1. Alle Menschen sind sterblich.
 2. Sabine ist ein Mensch.
- Also: Sabine ist sterblich.

Schlüssigkeit ist eine hinreichende Bedingung für Korrektheit.

Bei einem zirkulären Schluss liefert das Argument keinen unabhängigen Grund für die Konklusion: Um die Wahrheit der Prämisse zu prüfen, muss man die Wahrheit der Konklusion kennen. Mit anderen Worten: Die Konklusion ist (meist versteckter) Teil der Prämissen.

- Bochum liegt im Ruhrgebiet und hat über 300000 Einwohner.
- Also: Bochum liegt im Ruhrgebiet.

▷ Logische Form von Argumenten

Betrachten wir die folgenden beiden Beispiele:

1. Wenn Hans der Mörder ist, war er am Tatort.
 2. Hans war nicht am Tatort.
- Also: Hans ist nicht der Mörder.
-
1. Wenn Anna die Siegerin ist, hat sie am Wettbewerb teilgenommen.
 2. Anna hat nicht am Wettbewerb teilgenommen.
- Also: Anna ist nicht die Siegerin.

Beide Argumente haben die gleiche Struktur: Sie entstehen aus derselben **Argumentform (logische Form)**:

1. Wenn P , dann Q
 2. $\neg Q$
- Also: $\neg P$

Man sagt, die Argumente seien Einsetzungsinstanzen derselben Argumentform.

Ein Argument ist deduktiv korrekt, wenn alle strukturgleichen Argumente mit wahren Prämissen auch wahre Konklusionen haben. Es ist falsch, wenn es ein Gegenbeispiel gibt.

Das nicht korrekte Argument

1. Kein Papagei ist ein Säugetier.
 2. Kein Säugetier ist ein Fisch.
- Also: Kein Papagei ist ein Fisch.

ist eine Einsetzungsinstanz der nicht korrekten Argumentform

1. Kein P ist ein Q .
 2. Kein Q ist ein R .
- Also: Kein P ist ein R .

Hierbei kann man die Formulierung „Kein P ist ein Q “ noch übersetzen in „Für alle $x \in P$ gilt: $x \notin Q$.“ (Man mache sich an z.B. an Venn-Diagrammen klar, dass die obige Argumentform falsch ist!)

Weitere Beispiele für (korrekte) Formen:

1. 1. Wenn P , dann Q (kurz: $P \Rightarrow Q$)
 2. P
 Also: Q
2. Entweder P oder Q .
 $\neg P$
 Also: Q
3. Wenn P , dann Q
 $\neg Q$
 Also: P

Zum Abschluss wollen wir noch die logische Form des Darwin-Zitats angeben. Hierzu führen wir die folgenden Abkürzungen ein.

- A : Das evolutionäre Ziel ist die **A**rterhaltung.
 L : Der **L**ebensraum auf der Erde ist begrenzt.
 G : Es gibt einen **G**eburtenüberschuss.
 B : Auswahl der **B**esten.
 U : Art wächst **u**nbeschränkt über alle Grenzen.

Die logische Form ist dann:

1. A
 2. L
 3. G
 4. $G \wedge \neg B \Rightarrow U \quad \Leftrightarrow \quad \neg U \Rightarrow \neg G \vee B$
 5. $L \wedge U \Rightarrow \neg A \quad \Leftrightarrow \quad A \Rightarrow \neg L \vee \neg U$
- Also: B

Die Implikationen sind hierbei äquivalent umgeformt.

▷ Formale Überprüfung auf Korrektheit

Ein Argument ist formal korrekt, wenn die Wahrheit der Prämissen die Wahrheit der Konklusion logisch erzwingt. Um dies zu zeigen, kann man eine logische Schlusskette angeben, an deren Ende die Konklusion steht. Für das Darwin-Zitat kann man etwa folgende Schlusskette angeben:

$$\left. \begin{array}{l} A \Rightarrow \neg L \vee \neg U \\ L \\ G \end{array} \right\} \Rightarrow \neg U \Rightarrow \neg G \vee B \left. \right\} \Rightarrow B$$

Zu lesen: Aus der Prämisse A folgt $\neg L$ oder $\neg U$. Zusammen mit der Prämisse L folgt die Gültigkeit von $\neg U$ (da ja $\neg L = F$). Aus $\neg U$ folgt jedoch $\neg G$ oder B . Da G wahr ist, ist $\neg G$ falsch, also muss B wahr sein. Logisch!

Die formale Inkorrektheit kann man zeigen, indem man ein Gegenbeispiel angibt.

Besteht ein Argument aus vielen Prämissen und Teilargumenten, so kann es sehr schwer sein, eine Herleitungskette anzugeben. Man kann jedoch die folgende sog. direkte Methode verwenden. Die direkte Methode wendet die Wahrheitstafelmethode auf das Argument an. Man stellt eine Wahrheitstafel auf, in der alle möglichen Belegungen der Aussagenvariablen mit Wahrheitswerten verzeichnet sind. Für alle Prämissen und die Konklusion stellt man den zugehörigen Wahrheitswert fest. Die Argumentform ist korrekt, wenn in den Zeilen, in denen alle Prämissen den Wahrheitswert W haben, auch die Konklusion den Wahrheitswert W hat. Ansonsten ist das Argument nicht korrekt.

Für das Darwin-Zitat müssen wir lediglich die Tafel betrachten, bei der B und U frei variieren, da ja A , L und G unmittelbar als Prämissen auftreten.

B	U	$(G \wedge \neg B \Rightarrow U)$	$\wedge$	$(L \wedge U \Rightarrow \neg A)$	$\Rightarrow$	B
W	W	W	F	F	W	W
W	F	W	W	W	W	W
F	W	W	F	F	W	F
F	F	F	F	W	W	F

Der Ergebnis dieser systematischen Fleißarbeit ist ein wasserdichter Beweis der Korrektheit des Arguments.

1.3

1.3 Zahlssysteme und elementares Rechnen

In den Naturwissenschaften gehört der Umgang mit Anzahlen (z.B. Zählen von Sozialkontakten bei Tierbeobachtungen), Verhältniszahlen (z.B. Ansetzen einer 70 %-igen Alkohol-Lösung im Labor), sowie Messwerte „mit Nachkommastellen“ (etwa Gewichtsangaben) zum täglichen Brot. Die mathematischen Entsprechungen sind die natürlichen Zahlen $\mathbb{N}$, die rationalen Zahlen $\mathbb{Q}$ sowie die reellen Zahlen $\mathbb{R}$. Im Folgenden werden einige wichtige Sachverhalte und Rechenregeln zusammengestellt.

► 1.3.1 Die natürlichen Zahlen

Die beim Zählen von Dingen auftretenden **natürlichen Zahlen** werden mit

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

bezeichnet. Nimmt man die 0 hinzu, so schreibt man $\mathbb{N}_0 = \{0, 1, 2, \dots\}$.

Mit natürlichen Zahlen rechnet man, wie man es aus der Schule kennt. Man kann sie addieren und multiplizieren, wobei Punkt- vor Strich-Rechnung geht. Also ist $2 \cdot 3 + 5 = 11$ und nicht 16.

▷ Division mit Rest

Für natürliche Zahlen ist die Division mit Rest, DIV, erklärt. Um $a \text{ DIV } b$ zu ermitteln, sucht man ein $k \in \mathbb{N}$, so dass $a = k \cdot b + r$ mit einem Rest r für den gilt: $0 \leq r < b$. Also ist $17 \text{ DIV } 5 = 3 \text{ REST } 2$, da $17 = 3 \cdot 5 + 2$.

b heißt **Teiler** von a , wenn sich kein Rest ergibt, d.h., wenn a ein Vielfaches von b ist:

$$a = k \cdot b$$

für ein $k \in \mathbb{N}$. Natürlich kann man immer schreiben $a = k \cdot b$, wenn man $k = 1$ und $b = a$ setzt oder $k = a$ und $b = 1$, aber dies sind uninteressante Fälle. b ist dann kein **echter Teiler**. Zahlen, die keine echten Teiler besitzen heißen **Primzahlen**:

$$2, 3, 5, 7, 11, 13, 17, 19, \dots$$

Primzahlen sind also nur durch 1 und durch sich selbst teilbar.

1.3.1

Anmerkung 1.3.1 Mit dem *Sieb des Eratosthenes* kann man alle Primzahlen ermitteln, die kleiner als eine vorgegebene Zahl n sind. Man schreibe alle natürlichen Zahlen von 2 bis n auf. Nun streiche man die 2 und jede zweite auf 2 folgende Zahl.

Ist p die erste nicht gestrichene Zahl, so markiere man diese und streiche jede p -te darauf folgende Zahl.

▷ Primfaktorzerlegung

Ein fundamentales Resultat der Mathematik besagt, dass man jede natürliche Zahl ≥ 2 in ein Produkt von Primfaktoren zerlegen kann. Die Primfaktorzerlegung wird verwendet, um Brüche zu kürzen. Zudem spielt sie eine wichtige Rolle bei der Chiffrierung und Dechiffrierung von Texten.

Ein Primfaktor ist eine Potenz einer Primzahl p , also von der Form p^q mit $q \in \mathbb{N}$. Der Satz von der Primfaktorzerlegung besagt nun, dass es zu jeder natürlichen Zahl $a \geq 2$ endlich viele Primzahlen $p_1, \dots, p_n$ gibt mit zugehörigen Exponenten $r_1, \dots, r_n$, so dass

$$a = p_1^{r_1} \cdot \dots \cdot p_n^{r_n} = \prod_{i=1}^n p_i^{r_i}.$$

Die Primfaktorzerlegungen der Zahlen 2 bis 10 lauten:

$$2 = 2^1, 3 = 3^1, 4 = 2^2, 5 = 5^1, 6 = 2 \cdot 3, 7 = 7^1, 8 = 2^3, 9 = 3^2, 10 = 2 \cdot 5.$$

Dann geht es weiter mit

$$11 = 11^1, 12 = 3 \cdot 2^2, 13 = 13^1, 14 = 2 \cdot 7, 15 = 3 \cdot 5$$

und

$$16 = 2^4, 17 = 17^1, 18 = 2 \cdot 3^2, 19 = 19^1, 20 = 2^2 \cdot 5.$$

➤ 1.3.2 Die ganzen Zahlen

Die ganzen Zahlen erhält man aus den natürlichen Zahlen durch Hinzunahme der 0 und aller negativen Zahlen.

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

Die Operationen $+$, $\cdot$, und Division mit Rest führen nicht aus dem Bereich der ganzen Zahlen hinaus. Man sagt auch: $\mathbb{Z}$ ist abgeschlossen bzgl. dieser Operationen.

➤ 1.3.3 Die rationalen Zahlen (Bruchzahlen)

Bruchzahlen treten in natürlicher Weise bei der Angabe von Verhältnissen auf: „Um die Substanz A anzusetzen, mische man 3 Teile der Flüssigkeit B und 5 Teile der Flüssigkeit C .“ Insgesamt hat man dann 8 Teile (genauer:

Volumen- oder Gewichtseinheiten), so dass die Mischung zu $\frac{3}{8}$ aus B und zu $\frac{5}{8}$ aus C besteht.

Die rationalen Zahlen bestehen aus allen Bruchzahlen:

$$\mathbb{Q} = \left\{ \frac{p}{q} : p, q \in \mathbb{Z}, q \neq 0 \right\}.$$

Bei einem Bruch $\frac{p}{q}$ heißt p **Zähler** und q **Nenner**. Bruchzahlen sollten immer in gekürzter Form angegeben werden. Einen Bruch kann man kürzen, wenn sowohl im Zähler als auch im Nenner ein gemeinsamer Faktor steht:

$$\frac{a \cdot x}{b \cdot x} = \frac{a}{b}.$$

Bei einem gekürzten Bruch haben Zähler und Nenner keinen gemeinsamen Faktor. Gibt man einen Bruch nicht in gekürzter Form an, so sollte das inhaltlich begründet sein, bspw. weil man das Mischungsverhältnis in praktikabler Form angeben will.

Kürzen eines Bruchs: Häufig sieht man sofort, wie man einen Bruch kürzen kann: 12 durch 4 ist eben 3, also ist $\frac{12}{4} = 3$, aber es gibt auch ein *systematisches Verfahren*: Um einen Bruch $\frac{a}{b}$ zu kürzen, bildet man die Primfaktorzerlegung von Nenner und Zähler. Treten Primzahlen sowohl im Zähler als auch im Nenner auf, so kann man diese kürzen.

Rechenregeln für Brüche: Brüche werden multipliziert, indem man Zähler und Nenner einzeln multipliziert: Für alle $a, b, c, d \in \mathbb{N}$ mit $b, d \neq 0$ gilt

$$\frac{a}{b} \cdot \frac{c}{d} = \frac{ac}{bd}.$$

Der Bruch $\frac{a}{b}$ wird durch $\frac{c}{d}$ dividiert, indem man $\frac{a}{b}$ mit dem Kehrbuch $\frac{d}{c}$ multipliziert:

$$\frac{a}{b} : \frac{c}{d} = \frac{a}{b} \cdot \frac{d}{c} = \frac{ad}{bc}$$

$\frac{a}{b} : \frac{c}{d}$ ist hierbei eine andere Schreibweise für den entsprechenden Doppelbruch. Also:

$$\frac{\frac{a}{b}}{\frac{c}{d}} = \frac{a}{b} : \frac{c}{d} = \frac{ad}{bc}$$

► 1.3.4 Die reellen Zahlen

Die reellen Zahlen $\mathbb{R}$ kann man sich vorstellen als die Menge aller Punkte der unendlichen Zahlengeraden. Es stellt sich die Frage, ob die reellen Zahlen nicht dasselbe sind wie die rationalen Zahlen. Dies ist nicht der Fall: Es gibt „Lücken“ in $\mathbb{Q}$. Auf diese Lücken stößt man bereits, wenn man Wurzeln

betrachtet. Die positive Lösung der Gleichung

$$x^2 = 2$$

bezeichnet man mit $\sqrt{2}$. In anderen Worten: $\sqrt{2}$ ist diejenige positive Zahl, die ins Quadrat erhoben 2 ergibt. $\sqrt{2}$ kann aber nicht als Bruch geschrieben werden. Solche Zahlen heißen irrationale Zahlen.

Der Beweis dieser Aussage ist ein Paradebeispiel für einen indirekten Beweis. Bei einem indirekten Beweis wird nicht direkt die Behauptung nachgewiesen, sondern man führt die Negation der Behauptung zum Widerspruch. Indirekte Beweisführungen werden auch oft in der umgangssprachlichen Argumentation verwendet: „Mal angenommen, Sie hätten mit Ihrer Behauptung ... Recht. Dann ergibt sich doch wohl ..., was jedoch offenkundig falsch ist. Deshalb ist Ihre Behauptung nicht richtig!“

Angenommen, $\sqrt{2}$ wäre ein Bruch. Dann können wir den Bruch in gekürzter Form schreiben:

$$\sqrt{2} = \frac{p}{q},$$

wobei p und $q \neq 0$ keinen gemeinsamen Teiler haben. Wir werden im Folgenden zeigen, dass aus dieser Annahme folgt, dass p und q den gemeinsamen Teiler 2 haben. Dies liefert einen Widerspruch zur Annahme, dass $\sqrt{2}$ ein Bruch ist. Quadrieren von $\sqrt{2} = \frac{p}{q}$ liefert die Darstellung $2 = \frac{p^2}{q^2}$. Dann können wir $p^2 = 2q^2$ schreiben. Somit ist p^2 eine gerade Zahl. Ist dann auch p gerade? Angenommen, nein. Dann ist p ungerade, d.h. wir können schreiben: $p = 2 \cdot s + 1$ mit $s \in \mathbb{N}$. Ausmultiplizieren liefert dann $p^2 = (2 \cdot s + 1)^2 = 4 \cdot s^2 + 4 \cdot s + 1$. Also ist p^2 ungerade, was nicht sein kann. Also ist mit p^2 auch p gerade. Hieraus folgt nun: $q^2 = p^2/2 = (p/2) \cdot (p/2)$. Also ist q^2 eine gerade Zahl, und damit ist auch q eine gerade Zahl. Folglich haben p und q den gemeinsamen Teiler 2 im Widerspruch zur Annahme, dass der Bruch schon gekürzt ist.

▷ Summen- und Produkte

An vielen Stellen in diesem Text wird uns folgende Situation begegnen: Gegeben sind $x_1, \dots, x_n \in \mathbb{R}$ (als Platzhalter für konkrete Zahlen), die summiert oder multipliziert werden sollen. Die folgenden Kurzschreibweisen haben sich eingebürgert:

$$\sum_{i=1}^n x_i = x_1 + \dots + x_n, \quad \prod_{i=1}^n x_i.$$

i heißt hierbei *Laufvariable*.

1.4 1.4 Potenzen, Wurzeln

► 1.4.1 Motivation

Wachstum von Bakterien: Eine Bakterienkultur in Nährlösung wachse ausgehend von einem Populationsumfang $B_0 > 0$ pro Periode um den Faktor x . Dann liegen nach einer Periode $B_1 = B_0 \cdot x$ und nach zwei Perioden

$$B_2 = B_1 \cdot x = (x \cdot x) \cdot B_0 = x^2 \cdot B_0$$

Bakterien vor. Allgemein gilt für den Bestand nach n Perioden:

$$B_n = x^n B_0.$$

Den n -Perioden-Faktor erhält man also durch **Potenzieren**.

Ist umgekehrt der 2-Perioden-Faktor c mit $B_2 = cB_0$ bekannt, so gilt:

$$x^2 = c.$$

Es ist also eine **quadratische Gleichung** zu lösen:

$$x_1 = \sqrt{c} \quad \text{und} \quad x_2 = -\sqrt{c}$$

sind die beiden Lösungen, wobei hier nur x_1 biologisch relevant ist. Kennt man allgemein den Faktor c mit $B_n = cB_0$, so gilt: $x^n = c$.

► 1.4.2 Potenzen

Für $a \in \mathbb{R}$ und $n \in \mathbb{N}$ heißt

$$a^n = \underbrace{a \cdot \dots \cdot a}_n, \quad a^0 = 1,$$

n -te Potenz von a . a heißt **Basis** und n **Exponent**. Es gilt dann die rekursive Darstellung

$$a^n = a^{n-1} \cdot a, \quad n \geq 1.$$

Ferner setzt man

$$a^{-n} = \frac{1}{a^n}, \quad \text{falls } a \neq 0.$$

Rechenregeln: Für $p, q \in \mathbb{Q}$ gelten die Formeln:

$$\frac{a^p}{a^q} = a^{p-q}, \quad a^p \cdot a^q = a^{p+q}, \quad a^p \cdot b^p = (ab)^p,$$

$$\frac{a^p}{b^p} = \left(\frac{a}{b}\right)^p, \quad (a^p)^q = a^{pq}.$$

Achtung: $2^{4^3} = 2^{(4^3)} = 2^{64} \neq (2^4)^3 = 2^{12}$.

➤ 1.4.3 Wurzeln

Für $a > 0$ und $n \in \mathbb{N}$ heißt $b > 0$ mit

$$b^n = a$$

n -te **Wurzel von** a . Schreibweise: $b = \sqrt[n]{a} = a^{1/n}$. a heißt **Radikand**. Also:

$$b = \sqrt[n]{a} = a^{\frac{1}{n}} \Leftrightarrow b^n = a.$$

Ist $a > 0$ und $r = \frac{p}{q} \in \mathbb{Q}$ ein Bruch mit $q > 0$, so ist

$$a^r = a^{\frac{p}{q}} = (a^{\frac{1}{q}})^p = (\sqrt[q]{a})^p.$$

Gerade Wurzeln ($n = 2, 4, 6, \dots$) sind nur für positive Zahlen definiert, *ungerade* können auch für negative definiert werden. So ist $\sqrt[3]{-8} = -2$, da $(-2) \cdot (-2) \cdot (-2) = -8$.

Rechenregeln: Da Wurzeln Potenzen mit rationalen Exponenten sind, übertragen sich die Rechenregeln.

Beispiel 1.4.1 Hier einige Beispiele:

1.4.1

1. $\sqrt[3]{27} = 27^{1/3} = 3$, da $3 \cdot 3 \cdot 3 = 27$.
2. $\frac{\sqrt{20}}{\sqrt{5}} = \sqrt{\frac{20}{5}} = \sqrt{4} = 2$.
3. $\sqrt[3]{-54} = \sqrt[3]{(-1) \cdot 2 \cdot 27} = \sqrt[3]{-1} \cdot \sqrt[3]{2} \sqrt[3]{27} = (-1)\sqrt{2}3 = -3\sqrt{2}$.
4. $\frac{\sqrt{x^2-y^2}}{\sqrt{x-y}} = \sqrt{\frac{(x+y)(x-y)}{x-y}} = \sqrt{x+y}$.

Die Wurzel aus einer Zahl ist stets eindeutig bestimmt. So ist $\sqrt{4} = 2$ und nicht ± 2 . Die zugehörige Gleichung $x^2 = 4$ hat jedoch *zwei* Lösungen $x_1 = \sqrt{4} = 2$ und $x_2 = -\sqrt{4} = -2$.

➤ 1.4.4 Lösen von Potenzgleichungen

Gleichungen, in denen Potenzen vorkommen, heißen Potenzgleichungen. Bei der mathematischen Behandlung naturwissenschaftlicher Phänomene stößt man sehr rasch auf solche Gleichungen. Besonders wichtig sind hierbei Polynomgleichungen.

Ein **Polynom** in der Variablen x ist ein Ausdruck der Form

$$a_0 + a_1x + a_2 \cdot x^2 + \dots + a_p \cdot x^p.$$

Hierbei sind $a_0, \dots, a_p$ reelle Zahlen, genannt **Koeffizienten**. Beispiele:

$$2 + 3x, \quad 2x^2 + 5x - 3, \quad x^4 - 3x^3 + x.$$

Es tauchen also Potenzen von x auf, die jeweils mit Koeffizienten multipliziert und dann aufsummiert werden.

▷ Quadratische Gleichungen

Wir werden an vielen Stellen auf *quadratische Gleichungen* stoßen. Das sind Gleichungen der Form

$$ax^2 + bx + c = 0$$

mit $a \neq 0$ und $b, c \in \mathbb{R}$. Gesucht werden also Nullstellen x des quadratischen Polynoms $p(x) = ax^2 + bx + c$. Man überführt die Gleichung zunächst in die Normalform und führt dann eine quadratische Ergänzung durch. Man spricht hierbei von Normalform, wenn $a = 1$ ist. Mit $p = b/a$, $q = c/a$ erhält man:

$$\begin{aligned} x^2 + px + q &= 0 \\ \Leftrightarrow x^2 + px + (p/2)^2 &= (p/2)^2 - q \\ \Leftrightarrow (x + p/2)^2 &= (p/2)^2 - q \\ \Leftrightarrow x + p/2 &= \pm \sqrt{(p/2)^2 - q}. \end{aligned}$$

Der letzte Schritt ist korrekt, wenn der Radikand

$$D = (p/2)^2 - q$$

nicht negativ ist. Somit erhalten wir die bekannte Lösungsformeln:

$$x_1 = -\frac{p}{2} - \sqrt{\left(\frac{p}{2}\right)^2 - q}, \quad x_2 = -\frac{p}{2} + \sqrt{\left(\frac{p}{2}\right)^2 - q}.$$

Für $D < 0$ gibt es keine Lösung, für $D = 0$ genau eine Lösung und für $D > 0$ gibt es zwei Lösungen. D unterscheidet (diskriminiert) zwischen den verschiedenen Lösungstypen. D heißt daher **Diskriminante**.

Gelegentlich ist der **Satz von Vieta** nützlich: Zwischen den Lösungen (Nullstellen) und den Koeffizienten gilt der Zusammenhang

$$p = -(x_1 + x_2), \quad q = x_1 \cdot x_2.$$

1.4.1

Anmerkung 1.4.1 Potenzgleichungen vom Grad $n \geq 3$ sind i.a. nicht explizit bzw. nicht vollständig lösbar. Für gewisse Sonderfälle gibt es jedoch spezielle Lösungsmethoden. Eine Gleichung der Form $ax^4 + bx^2 + c = 0$ kann man etwa durch die Substitution $z = x^2$ auf eine quadratische Gleichung in z zurückführen.

▷ **Wurzelgleichungen**

Dies sind Gleichungen, bei denen die Variable im Radikand steht. Etwa:

$$\sqrt{x-1} + 3 = x \quad \text{oder} \quad (x^2 - 1)^{1/3} = 0.$$

Hier ist zunächst der Definitionsbereich zu bestimmen. Man versucht dann Lösungen durch Potenzieren zu finden. Hierdurch können neue Lösungen hinzukommen! (Warum?) Man muss also testen, ob die Lösungen der durch Potenzieren gefundenen Gleichungen auch Lösungen der Ausgangsgleichung sind.

Beispiel 1.4.2 Der Definitionsbereich der Gleichung

1.4.2

$$\sqrt{x-1} + 3 = x$$

ist $D = \{x \in \mathbb{R} | x \geq 1\}$. Isoliere nun die Wurzel und quadriere beide Seiten. Für $x \geq 1$ gilt:

$$\begin{aligned} \sqrt{x-1} + 3 &= x \\ \Leftrightarrow \sqrt{x-1} &= x-3 \\ \Rightarrow x-1 &= (x-3)^2 = x^2 - 6x + 9 \\ \Leftrightarrow x &\in \{2, 5\}. \end{aligned}$$

Von den beiden Lösungen der quadratischen Gleichung ist nur $x = 5$ Lösung der Ausgangsgleichung.

➤ **1.4.5 Prozentrechnung, Rechnen mit Wachstumsraten**▷ **Begriffsbildungen**

Bei empirischen Untersuchungen hat man es häufig mit zeitlich geordneten *Bestandsgrößen* zu tun. Hierunter fallen Populationsumfänge und ganz allgemein Zählungen von Dingen, aber auch Messungen von Größen wie Volumina oder Gewichte. Wir wollen davon ausgehen, dass solch eine *Zeitreihe* $B_0, B_1, \dots, B_n$ vorliegt, wobei B_0 den Ausgangsbestand im Zeitpunkt t_0 bezeichnet, und B_i den Wert am Ende der i -ten Periode $(t_{i-1}, t_i]$. Dann heißt

$$w_i = B_i / B_{i-1} \quad \Leftrightarrow \quad B_i = B_{i-1} \cdot w_i$$

Wachstumsfaktor (engl: *growth factor*) der i -ten Periode. Der Wachstumsfaktor w_i ist also derjenige Faktor, mit dem der Wert B_{i-1} der Vorperiode zu multiplizieren ist, um den Wert der aktuellen Periode B_i zu erhalten. Die

zugehörige **Wachstumsrate** (engl: *growth rate*) ist definiert als

$$r_i = w_i - 1 \quad \Leftrightarrow \quad B_i = (1 + r_i)B_{i-1}.$$

$100 \cdot r_i \%$ ist also der prozentuale Zuwachs bzw. die prozentuale Schrumpfung während der i -ten Periode. Formal ist das Prozentzeichen als multiplikativer Faktor $\% = \frac{1}{100}$ definiert. Ist $r_i = 0.05$ so entspricht dies einem Wachstum von 5%.

Der Bestand B_n am Ende des Betrachtungszeitraums berechnet sich aus B_0 und den Wachstumsfaktoren bzw. Wachstumsraten durch

$$\begin{aligned} B_n &= B_{n-1} \cdot w_n \\ &= B_{n-2} \cdot w_{n-1} w_n \\ &\vdots \\ &= B_0 \cdot w_1 \cdot \dots \cdot w_n \\ &= B_0 \prod_{i=1}^n w_i. \end{aligned}$$

Setzt man $w_i = 1 + r_i$ ein, so erhält man die Formel

$$B_n = B_0 \cdot (1 + r_1) \cdot \dots \cdot (1 + r_n) = B_0 \prod_{i=1}^n (1 + r_i).$$

▷ Durchschnittlicher Wachstumsfaktor

Der durchschnittliche Wachstumsfaktor w^* ist definiert als derjenige Wachstumsfaktor, der bei Anwendung in allen n Perioden zum (vorgegebenen) Endbestand B_n führt. D.h.:

$$B_n = B_0 \underbrace{w^* \cdot \dots \cdot w^*}_n = B_0 \cdot (w^*)^n.$$

Division durch B_0 liefert:

$$(w^*)^n = w_1 \cdot \dots \cdot w_n \Leftrightarrow w^* = \sqrt[n]{w_1 \cdot \dots \cdot w_n}.$$

w^* ist also durch das *geometrische Mittel* der n Wachstumsfaktoren $w_1, \dots, w_n$ gegeben.

Abschätzung des Wachstumsfaktors

Eine einfache untere Abschätzung des Wachstumsfaktors für n Perioden erhält man durch die **Bernoulli'sche Ungleichung**: Sei $x \geq -1$. Dann gilt für alle $n \in \mathbb{N}$:

$$(1 + x)^n \geq 1 + n \cdot x.$$

Die rechte Seite kann man leicht im Kopf ausrechnen.

Anwendung: Eine Population wachse ausgehend von B_0 um $x \cdot 100\%$ pro Periode. Der Endbestand $B_n = B_0 \cdot (1+x)^n$ beträgt dann mindestens $B_0 \cdot (1+n \cdot x)$.

▷ **Durchschnittliche Wachstumsrate**

Die durchschnittliche Wachstumsrate w^* ist definiert als diejenige Wachstumsrate, die bei Anwendung in allen n Perioden zum vorgegebenen Endbestand B_n führt. Einsetzen von $w^* = 1 + r^*$ und $w_i = 1 + r_i$ liefert daher:

$$r^* = \sqrt[n]{(1+r_1) \cdot \dots \cdot (1+r_n)} - 1$$

1.5 Kombinatorik

1.5

Bei der Planung eines Experiments steht man mitunter vor dem Problem, aus einer großen Grundgesamtheit von n Objekten eine Teilauswahl (Stichprobe) auszuwählen, da es unmöglich ist, das Experiment für alle Elemente der Grundgesamtheit durchzuführen. Man kann sich die n Objekte als n durchnummerierte Kugeln denken, die in einer Urne liegen (Urnenmodell).

Stichprobe mit/ohne Zurücklegen: Je nachdem, ob das ausgewählte Objekt wieder zurückgelegt wird oder nicht, spricht man von einer Auswahl (Stichprobe) mit bzw. ohne Zurücklegen.

Stichprobe in/ohne Reihenfolge: Man spricht von einer Stichprobe in Reihenfolge, wenn es auf die Reihenfolge der Züge ankommt. Das Ergebnis wird dann durch ein k -Tupel $(\omega_1, \dots, \omega_k)$ beschrieben, wobei ω_i das Ergebnis des i -ten Zuges ist. Kommt es hingegen nicht auf die Reihenfolge an, so spricht man von einer Stichprobe ohne Reihenfolge. Sind Mehrfachziehungen ausgeschlossen, so können wir das Ergebnis als Menge statt als Vektor aufschreiben: $\{\omega_1, \dots, \omega_k\}$. Sind Mehrfachziehungen möglich, so ist relevant, wie oft jedes Objekt ausgewählt wurde. Das Ergebnis eines Zuges wird daher durch ein n -Tupel $(\omega_1, \dots, \omega_n)$ beschrieben, wobei nun ω_i angibt, wie oft das Objekt i ausgewählt wurde. Da k - mal gezogen wird, ist die Summe der ω_i gerade k .

Durch diese beiden Charakterisierungen ergeben sich vier verschiedene Urnenmodelle. Entscheidend ist nun zu untersuchen, wie viele verschiedene Möglichkeiten es gibt, k Objekte auszuwählen.

▷ **Modell I: Stichprobe in Reihenfolge und mit Zurücklegen**

Die Menge aller möglichen Stichproben ist durch

$$\Omega_I = \{\omega = (\omega_1, \dots, \omega_k) : \omega_i \in \{1, \dots, n\}, i = 1, \dots, k\}$$

beschrieben, wobei ω_i die Nummer des i -ten ausgewählten Objektes ist. Da jedes gezogene Objekt wieder zurückgelegt wird, gibt es bei jedem Zug genau n verschiedene Möglichkeiten. Insgesamt gibt es also

$$|\Omega_I| = \underbrace{n \cdot \dots \cdot n}_k = n^k$$

verschiedene mögliche Stichproben.

▷ **Modell II: Stichprobe in Reihenfolge und ohne Zurücklegen**

Da man nicht zurücklegt, besteht jede Stichprobe $(\omega_1, \dots, \omega_k)$ aus k unterschiedlichen Objekten, d.h. für verschiedene Ziehungen sind die gezogenen Objekte verschieden. Kurz: Für $i \neq j$ gilt $\omega_i \neq \omega_j$.

$$\Omega_{II} = \{\omega = (\omega_1, \dots, \omega_k) : \omega_i \in \{1, \dots, n\}, \omega_i \neq \omega_j, \text{ für } i \neq j, 1 \leq i, j \leq k\}.$$

Wieviele Elemente hat Ω_{II} ? Beim 1. Zug gibt es n Objekte, die zur Auswahl stehen.

Beim 2. Zug gibt es $n - 1$ Objekte, die zur Auswahl stehen.

Beim 3. Zug gibt es $n - 2$ Objekte, die zur Auswahl stehen.

usw. Also gibt es bei k Zügen genau

$$|\Omega_{II}| = n_k := n(n-1) \cdot \dots \cdot (n-k+1)$$

Möglichkeiten, k Objekte in Anordnung aus n Objekten auszuwählen.

Für $k = n$ erhält man alle **Permutationen** der n Objekte. Die Anzahl der möglichen Permutationen

$$n! := n \cdot (n-1) \cdot \dots \cdot 2 \cdot 1.$$

heißt n **Fakultät**. Man setzt $0! = 1$.

▷ **Modell III: Stichprobe ohne Reihenfolge und ohne Zurücklegen**

Da keine Mehrfachziehungen möglich sind und es nicht auf die Reihenfolge ankommt, kann das Stichprobenergebnis als Menge geschrieben werden:

$$\Omega_{III} = \{\{\omega_1, \dots, \omega_k\} : \omega_i \in \{1, \dots, n\}, \omega_i \neq \omega_j \text{ für } i \neq j, 1 \leq i, j \leq k\}.$$

Um die Anzahl der Möglichkeiten, aus n Objekten k Objekte ohne Beachtung der Reihenfolge auszuwählen, gehen wir von folgender Überlegung aus: Wir können zunächst die Reihenfolge beachten (also das Ergebnis als Tupel aufschreiben) und dann beim Abzählen all diejenigen Stichproben nur jeweils einmal zählen, die durch eine Umordnung auseinander hervorgehen. So liefern etwa die sechs Tupel

$$(1, 2, 3), (1, 3, 2), (2, 1, 3), (2, 3, 1), (3, 1, 2), (3, 2, 1)$$

jeweils dieselbe Menge $\{1, 2, 3\}$, wenn man die Reihenfolge „vergisst“. Statt n_k Möglichkeiten (mit Reihenfolge) gibt es also nur

$$\binom{n}{k} = \frac{n_k}{k!} = \frac{n!}{(n-k)!k!}$$

Die Zahl $\binom{n}{k}$ heißt **Binomialkoeffizient n über k** . Es gilt:

$$\binom{n}{1} = \frac{n!}{1!n!} = 1, \quad \binom{n}{n} = \frac{n!}{n!(n-n)!} = 1$$

$\binom{n}{k}$ gibt also die Anzahl der Möglichkeiten an, aus einer Obermenge mit n Objekten eine k -elementige Teilmenge auszuwählen. Anders ausgedrückt: $\binom{n}{k}$ ist die Anzahl der Möglichkeiten, n Objekte auf zwei Klassen so aufzuteilen, dass sich in einer Klasse k Objekte und in der anderen Klasse $n-k$ Objekte befinden. Diese Zuordnung kann man durch Angabe einer Teilmenge der Zahlen $1, \dots, n$ darstellen. So bedeutet die Menge $\{1, 3\}$, dass die Objekte 1 und 3 in die eine Klasse kommen und die Objekte 2, 4, 5, $\dots, n$ in die andere Klasse. Diesen Zusammenhang und die Formel $\frac{n!}{(n-k)!k!}$ kann man sich auch folgendermaßen klar machen:

1. Schreibe alle n Objekte hintereinander:

$$1, 2, 3, \dots, n$$

Es gibt genau $n!$ verschiedene Permutation dieser Objekte.

2. Klammere die ersten k Stellen und die letzten $n-k$ Stellen ein. Für $(1, \dots, n)$ erhält man also:

$$\underbrace{(1, 2, 3, \dots, k)}_k \underbrace{(k+1, k+2, k+3, \dots, n)}_{n-k}$$

3. Durch Permutation der ersten k Elemente (untereinander!) und der letzten $n-k$ Stellen erhält man alle $k!(n-k)!$ Permutationen der Tupel, die zu derselben Menge führen. Also gibt es genau $\frac{n!}{k!(n-k)!}$ verschiedene Möglichkeiten:

$$|\Omega_{III}| = \frac{n!}{k!(n-k)!}$$

Zwei wichtige Eigenschaften des Binomialkoeffizienten:

1. $\binom{n}{k} = \binom{n}{n-k}$
Herleitung: in der obigen Herleitung sind die Rollen von k und $n-k$ vertauschbar.

$$2. \binom{n+1}{k+1} = \binom{n}{k} + \binom{n}{k+1}.$$

Herleitung: Jede ausgewählte Teilmenge mit $k+1$ Elementen aus der Menge $\{1, \dots, n+1\}$ lässt sich einem der beiden folgenden Fälle zuordnen. Können wir diese Fälle abzählen, so ergibt sich $\binom{n+1}{k+1}$ als Summe der beiden Anzahlen.

Fall 1: $n+1$ ist in der Auswahl nicht vorhanden. Dann stammen alle $k+1$ Elemente aus der Menge $\{1, \dots, n\}$, und es gibt genau $\binom{n}{k+1}$ Möglichkeiten, dies zu tun.

Fall 2: $n+1$ ist in der Auswahl enthalten. Dann stammen die restlichen k Elemente aus der Menge $\{1, \dots, n\}$, und es gibt genau $\binom{n}{k}$ Möglichkeiten, dies zu tun.

▷ Das Pascal'sche Dreieck

Die zweite Formel besagt, dass man Binomialkoeffizienten der Reihe nach berechnen kann. Schreibt man die Binomialkoeffizienten in ein Dreiecks-Schema, oben mit $\binom{1}{1}$ beginnend, so dass in der n -ten Zeile die Binomialkoeffizienten

$$1 = \binom{n}{1}, \binom{n}{2}, \dots, \binom{n}{n} = 1$$

stehen, so berechnet sich jeder Eintrag als Summe der beiden über ihm stehenden.

$$\binom{1}{0} = 1 \quad \binom{1}{1} = 1$$

$$\binom{2}{0} = 1 \quad \binom{2}{1} = 2 \quad \binom{2}{2} = 1$$

$$\binom{3}{0} = 1 \quad \binom{3}{1} = 3 \quad \binom{3}{2} = 3 \quad \binom{3}{3} = 1$$

$$\binom{4}{0} = 1 \quad \binom{4}{1} = 4 \quad \binom{4}{2} = 6 \quad \binom{4}{3} = 4 \quad \binom{4}{4} = 1$$

▷ Der binomische Lehrsatz

Der binomische Lehrsatz gibt an, wie man $(a+b)^n$ berechnet. Zunächst ist

$$\begin{aligned} (a+b)^0 &= 1 \\ (a+b)^1 &= a+b \\ (a+b)^2 &= a(a+b) + b(a+b) \\ &= a^2 + ab + ba + b^2 \\ &= a^2 + 2ab + b^2 \end{aligned}$$

Allgemein berechnet man das Produkt

$$(a+b)^n = (a+b) \cdot \dots \cdot (a+b),$$

indem man über alle Produkte summiert, die man erhält, wenn man aus jeder Klammer einen Term (entweder a oder b) auswählt. Nach Umordnen der Faktoren haben alle diese Produkte die Form

$$a^{n-i}b^i, \quad i = 0, \dots, n,$$

hängen also nur davon ab, wie oft ein a (bzw. b) ausgewählt wurde. Dies kann man auch so auffassen, dass von den n Klammern i ausgewählt werden, die ein b liefern, die anderen liefern dann ein a . Die Anzahl der Möglichkeiten, dies zu tun, ist gerade $\binom{n}{i}$. Somit erhalten wir die Formel:

$$(a+b)^n = \sum_{i=0}^n \binom{n}{i} a^{n-i} b^i.$$

Für $n = 3$:

$$\begin{aligned} (a+b)^3 &= \binom{3}{0} a^3 + \binom{3}{1} a^2 b + \binom{3}{2} a b^2 + \binom{3}{3} b^3 \\ &= a^3 + 3a^2 b + 3a b^2 + b^3. \end{aligned}$$

Die Koeffizienten liest man also aus dem Pascal'sche Dreieck ab.

▷ **Modell IV: Stichprobe ohne Reihenfolge mit Zurücklegen**

Hier sind Mehrfachziehungen möglich, aber die Reihenfolge, in der die Objekte ausgewählt werden, ist egal. Relevant ist nur noch, wie oft jede Kugel ausgewählt wird.

$$\Omega_{IV} = \{(\omega_1, \dots, \omega_n) : \omega_i \in \{0, 1, \dots, k\}, i = 1, \dots, n\}.$$

Um die Anzahl der möglichen Stichproben abzuzählen, überlegen wir uns, wie man die Ziehung praktisch durchführen kann: Man nehme ein Blatt Papier und ziehe $n - 1$ vertikale Trennstriche, so dass man n Felder erhält, die von 1 bis n durchnummeriert werden. Man wählt eine Kugel aus und notiert das Ergebnis durch einen Punkt in dem entsprechenden Feld. Die Kugel wird zurückgelegt. Am Ende hat man k kleine Punkte, die sich auf die n Felder verteilen.

Egal, wie genau unsere Stichprobe aussah, das Blatt Papier besteht in jedem Fall aus $n - 1 + k$ Objekten, nämlich $n - 1$ *Trennstrichen* und k *Punkten*, wobei sich die Anordnung aus der Stichprobe ergibt. Anders ausgedrückt: Jede Stichprobe ist eindeutig dadurch festgelegt, dass wir von den $n - 1 + k$ Objekten k als Punkte festlegen und die übrigen als Trennstriche. (Man

mache sich das an einigen Beispielen klar.) Es gibt genau

$$|\Omega_{IV}| = \binom{n-1+k}{k}$$

Möglichkeiten, diese Festlegung zu treffen.

▷ Zusammenfassung

Es gibt vier verschiedenen Arten, aus n Objekten k auszuwählen, je nachdem, ob es auf die Reihenfolge der Züge ankommt und ob Mehrfachziehungen zulässig sind.

Stichprobe vom Umfang k	mit Zurücklegen	ohne Zurücklegen
aus n Objekten $1, \dots, n$		

in Reihenfolge	$ \Omega_I = n^k$	$ \Omega_{II} = n_k$
ohne Reihenfolge	$ \Omega_{IV} = \binom{n+k-1}{k}$	$ \Omega_{III} = \binom{n}{k}$

▷ Multinomialkoeffizienten

Wir hatten gesehen, dass der Binomialkoeffizient $\binom{n}{k}$ die Anzahl der Möglichkeiten angibt, n verschiedene Objekte auf zwei Klassen so zu verteilen, dass sich gerade k in der ersten und $n - k$ in der zweiten Klasse befinden. Hat man die Objekte auf r Klassen so zu verteilen, dass sich in der i -ten Klasse gerade k_i Objekte befinden, so hat man gerade

$$\binom{n}{k_1 \dots k_r} = \frac{n!}{k_1! \dots k_r!}$$

Möglichkeiten, dies zu tun. Dieser Ausdruck heißt **Multinomialkoeffizient**. Die Herleitung verläuft analog wie oben: Nummeriere die Objekte von 1 bis n durch und klammere in jeder Permutation der Reihe nach k_1, k_2 , usw. bis schließlich k_n aufeinander folgende Zahlen ein. Die Zahlen in der i -ten Klammern werden der Klasse i zugeordnet. Für $(1, \dots, n)$ erhält man also:

$$\left(\underbrace{(1, \dots, k_1)}_{\rightarrow 1}, \underbrace{(k_1 + 1, \dots, k_1 + k_2)}_{\rightarrow 2}, \dots, \underbrace{(n - k_r + 1, \dots, n)}_{\rightarrow r} \right)$$

Die Zuordnung ändert sich nicht, wenn die Elemente in den Klammern untereinander permutiert werden. Die Anzahl $n!$ aller Permutationen muss daher durch $k_1! \dots k_r!$ dividiert werden.

1.6 Reelle Zahlenfolgen

1.6

► 1.6.1 Motivation

Beobachtet man das Wachstum einer Bakterienkultur unter Laborbedingungen bei begrenztem Nährstoff-Vorrat, so stellt man zunächst ein sehr rasches Wachstum fest, das allmählich nachläßt und schließlich zum Erliegen kommt. Bestimmt man - bspw. im Minuten-Takt - die Anzahl der Bakterien und bezeichnet die n -te Messung mit a_n , so erhält man eine aufsteigende Folge von Zahlen

$$a_1 \leq a_2 \leq a_3 \dots,$$

die sich einem oberen Wert a von unten anschmiegt. Wir werden später Modelle kennen lernen, die eine präzise Beschreibung des Wachstumsverhaltens ermöglichen. Hier wollen wir zunächst die empirische Beobachtung mathematisch präzisieren, dass sich eine Folge von reellen Zahlen an einen Wert a „anschmiegt“.

► 1.6.2 Begriffsbildung

Unter einer **Folge** (a_n) verstehen wir eine Menge von nummerierten reellen Zahlen $a_1, a_2, a_3, \dots$. Die Nummern, hier $1, 2, 3, \dots$, also die natürlichen Zahlen $\mathbb{N}$ heißen in diesem Kontext **Indizes** und die Gesamtheit aller Indizes heißt **Indexmenge**. a_k heißt das k -te **Folgenglied** mit Index k . Es ist üblich, die Nummerierung in der Form $1, 2, 3, \dots$ durchzuführen, je nach Anwendung verwendet man aber auch andere Indizes. Eine **endliche Folge** besteht lediglich aus endlich vielen Zahlen $a_1, \dots, a_N$, ansonsten spricht man von einer **unendlichen Folge**.

Folgen kann man auf zweierlei Weise anschaulich darstellen. Zunächst kann man in einem xy -Koordinatensystem die Indizes auf der x -Achse und die zugehörigen Folgenglieder auf der y -Achse abtragen. Alternativ hierzu kann man diese Punkte auf die x -Achse projizieren, also lediglich alle Folgenglieder auf dem reellen Zahlenstrahl markieren.

Das so gewonnene Bild einer Folge kann nun natürlich ganz verschieden aussehen. Es kann wirr sein oder Struktur besitzen. Insbesondere, wenn wir die zeitliche Entwicklung einer Population im Auge haben, a_n mithin die Populationsbestand am Ende der n -ten Periode bezeichnet, so ist es von besonderem Interesse zu erkennen, ob sich die Folge mit wachsendem n einem einzigen (endlichen) Wert a annähert. Dann sagt man, dass die Folge gegen den Grenzwert a konvergiert. Andernfalls divergiert die Folge. Nun muss es nicht so sein, dass ab einem bestimmten Index alle Folgenglieder mit a übereinstimmen. Wir wollen mit dem Konvergenzbegriff auch den Fall abdecken, dass sich die

Folge dem Grenzwert nur langsam *annähert*, und zwar egal ob von unten, oben oder hin und her springend. Ist dies der Fall, so können wir um den Grenzwert a ein beliebig kleines nicht leeres Intervall legen, und immer (d.h.: für alle solchen Intervalle) werden nur endlich viele Folgenglieder außerhalb des Intervalls liegen, aber unendlich viele innerhalb. Dies präzisieren wir in der folgenden Definition.

Eine Folge (a_n) heißt **konvergent**, wenn es eine Zahl $a \in \mathbb{R}$ gibt, so dass für jedes noch so kleine $\epsilon > 0$ alle Folgenglieder bis auf endlich viele im Intervall $(a - \epsilon, a + \epsilon)$ liegen.

D.h.: Für alle $\epsilon > 0$ gilt: Es gibt einen Index $n_0 \in \mathbb{N}$, so dass für alle $n \geq n_0$ gilt:

$$|a_n - a| < \epsilon$$

a heißt dann **Grenzwert** oder auch **Limes**. Man schreibt:

$$a = \lim_{n \rightarrow \infty} a_n \quad \text{oder} \quad a_n \rightarrow a, \text{ für } n \rightarrow \infty.$$

Eine Folge heißt **divergent**, wenn sie keinen Grenzwert besitzt. Konvergiert eine Folge (a_n) gegen 0, so spricht man von einer Nullfolge.

Es gilt

$$\lim_{n \rightarrow \infty} a_n = a \quad \Leftrightarrow \quad \lim_{n \rightarrow \infty} |a_n - a| = 0.$$

Der *Abstand* $|a_n - a|$ der Folgenglieder a_n von a strebt also genau dann für $n \rightarrow \infty$ gegen 0, wenn (a_n) konvergent mit Grenzwert a ist.

Aus dieser Umformulierung der Definition können wir ein erstes einfaches Konvergenzkriterium ableiten: Kann man den Abstand $|a_n - a|$ nach oben durch eine Nullfolge (b_n) abschätzen, so folgt $a_n \rightarrow a$, für $n \rightarrow \infty$.

1.6.1 Beispiel 1.6.1 Die Folge $a_n = 4 + 1/n$, $n \in \mathbb{N}$, konvergiert gegen $a = 4$. Denn für alle $n \in \mathbb{N}$ gilt:

$$|a_n - 4| = \frac{1}{n}$$

und $1/n$ ist eine Nullfolge.

1.6.2 Beispiel 1.6.2 Man spricht von einer **geometrischen Folge**, wenn

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

mit einem $q \in \mathbb{R}$ und einer Konstanten $c \in \mathbb{R}$.

Für $|q| < 1$ konvergiert (a_n) gegen 0.

Für $q = 1$ ist (a_n) konstant: $a_n = c$ für alle $n \in \mathbb{N}$.

Für $|q| > 1$ divergiert (a_n) .

Es gibt viele Kriterien, um eine Folge auf Konvergenz zu untersuchen. Wir wollen hier nur eines näher besprechen, das sehr anschaulich ist und einfach anzuwenden ist.

Eine wichtige Eigenschaft einer konvergenten Folge ist ihre Beschränktheit:

Eine Folge (a_n) ist **beschränkt**, wenn es eine Konstante K gibt, so dass:

$$|a_n| \leq K, \quad \text{für alle } n \in \mathbb{N}.$$

Ist nämlich eine Folge nicht beschränkt, so kann man keine Konstante finden, die alle Folgenglieder einfängt. Dann kann es aber keine Zahl a geben, in deren Nähe sich alle Folgenglieder ab einem gewissen Index aufhalten. Somit sind unbeschränkte Folgen nicht konvergent.

Konvergente Folgen sind also beschränkt, aber beschränkte Folgen nicht unbedingt konvergent. Man denke an Folgen, die sich periodisch verhalten.

Ist jedoch eine Folge beschränkt und **monoton wachsend**, d.h.

$$a_n \leq a_{n+1} \quad \text{für alle } n \in \mathbb{N},$$

so bleibt ihr nichts anderes übrig, als zu konvergieren: Mit wachsendem n werden die Werte a_n höchstens größer, sie können aber nicht beliebig groß werden. Solche Folgen konvergieren daher gegen die kleinste obere Schranke, die man finden kann. Genauso sind beschränkte und monoton fallende Folgen konvergent.

Der Folge

$$a_n = 4 + 1/n, \quad n \in \mathbb{N},$$

„sieht“ man den Grenzwert 4 direkt an. Es gibt jedoch auch konvergente Folgen, bei denen das nicht der Fall ist. Ein für die Biologie wichtiges Beispiel ist die **Euler'sche Zahl**

$$e = 2.71828.$$

(Leonhard Euler, 1707-1783). Es gilt:

$$e = \lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n.$$

Schaut man sich die Folge $a_n = (1 + 1/n)^n$ an, so glaubt man sofort, dass diese Folge beschränkt und streng monoton wachsend ist. Die entsprechenden Rechnungen sind jedoch recht diffizil.

Abbildung 1.1 illustriert einige konvergente Folgen, u.a. auch die Folge $(1 + 1/n)^n$. Abbildung 1.2 illustriert eine divergente Folge, die zunächst den Anschein erweckt, sie sei konvergent.

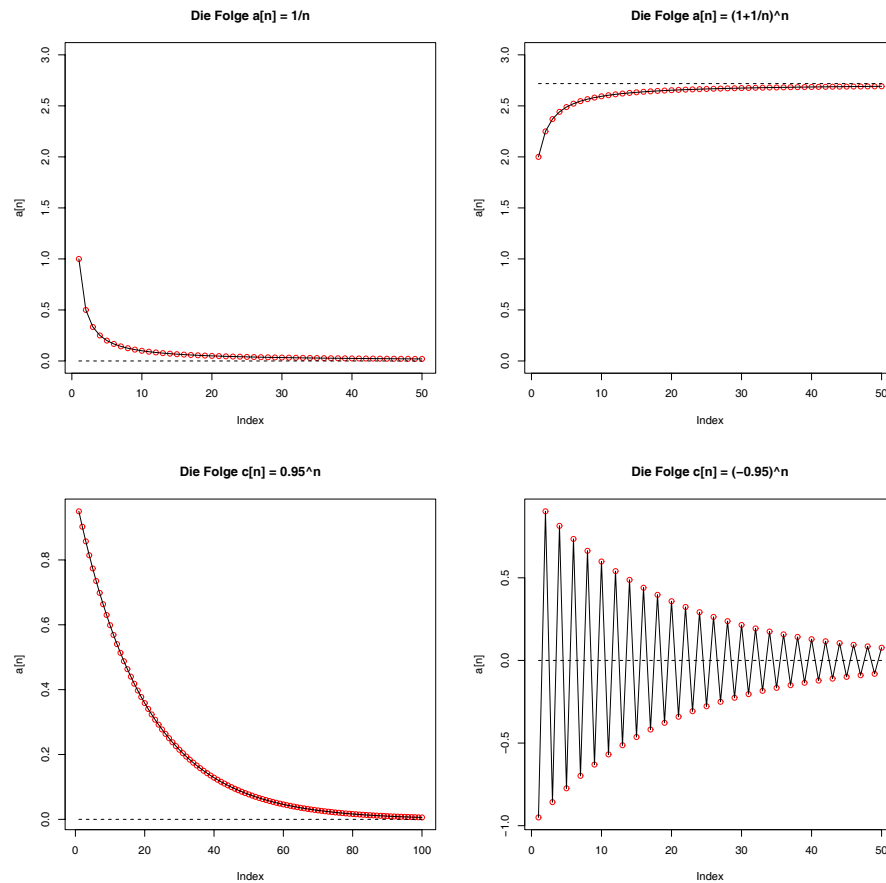


Abbildung 1.1: Einige konvergente Folgen. Die Folgenglieder sind durch Strecken verbunden, um die Abfolge besser zu veranschaulichen.

1.7 Reihen

► 1.7.1 Motivation

Einer Zelle werde zum Zeitpunkt $t = 0$ v_0 [ml] einer Substanz zugeführt. Bis zur Zeit $t = 1$ werden $p \cdot 100\%$ abgebaut. In $t = 1$ werden erneut v_0 [ml]

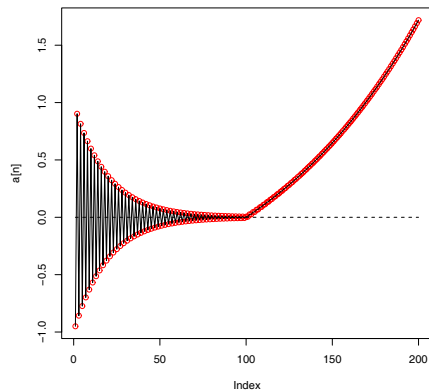


Abbildung 1.2. Eine divergente Folge

zugeführt. Dann befinden sich also

$$V_1 = v_0 \cdot q + v_0, \quad q = 1 - p,$$

Milliliter in der Zelle. Dieser Vorgang wird nun ad infinitum fortgeführt. Zur Zeit $t = 2$ befinden sich

$$V_2 = q(v_0 \cdot q + v_0) + v_0 = v_0 q^2 + v_0 q + v_0 \text{ [ml]}$$

in der Zelle. Zur Zeit $t = n$ erhält man den Ausdruck

$$V_n = v_0 \cdot q^n + v_0 \cdot q^{n-1} + \dots v_0 \cdot q + v_0.$$

Was passiert nun im Zeitablauf? Wächst die Folge (V_n) über alle Grenzen oder bleibt die Menge der Substanz beschränkt, da stets hinreichend viel abgebaut wird? In der obigen Formel für V_n können wir v_0 ausklammern:

$$V_n = v_0 \cdot (1 + q + q^2 + \dots q^n).$$

Wir erhalten die Antwort, wenn wir das Verhalten der Summe

$$S_n = 1 + q + q^2 + \dots + q^n = \sum_{i=0}^n q^i$$

für wachsendes n studieren.

► 1.7.2 Summen (Endliche Reihen)

Die Summe von endlich vielen Gliedern einer Folge heißt **endliche Reihe**.

$$(a_0, a_1, \dots, a_n) \rightarrow \sum_{i=0}^n a_i = a_0 + a_1 + \dots + a_n.$$

Für eine gegebene Folge hängt der Wert nur von n ab und kann für einige wichtige Spezialfälle explizit berechnet werden. Hier zwei Beispiele, die wir später verwenden werden.

1.7.1 Beispiel 1.7.1 Summe der ersten n Zahlen

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

Herleitung: Summiere die Zahlen 1 bis n zweimal:

$$\begin{array}{ccccccc} 1 & + & 2 & + & 3 & + & \dots & + & n \\ + & n & + & n-1 & + & n-2 & + & \dots & + & 1 & = & n(n-1) \end{array}$$

Die Summe der untereinander stehenden Zahlen ist jeweils $n+1$. Insgesamt ist die Summe der beiden Zeilen $n(n+1)$. Da wir jede Zahl doppelt gezählt haben, müssen wir das Ergebnis noch durch 2 dividieren.

1.7.2 Beispiel 1.7.2 Endliche arithmetische Reihe. Hier sind die Folgenglieder durch

$$a_n = a_0 + nd, \quad n = 0, 1, 2, \dots$$

gegeben. Die Folge startet in a_0 ; der Abstand zwischen den Folgengliedern ist stets d . Dann ist

$$\begin{aligned} \sum_{i=0}^n a_i &= a_0 + (a_0 + d) + (a_0 + 2d) + \dots + (a_0 + nd) \\ &= (n+1)a_0 + d(1 + 2 + 3 + \dots + n) \\ &= (n+1)a_0 + d \frac{n(n+1)}{2} \end{aligned}$$

► 1.7.3 Unendliche Reihen

Es sei (a_n) eine gegebene Folge reeller Zahlen und

$$S_n = a_0 + a_1 + \dots + a_n = \sum_{i=0}^n a_i$$

die n -te Partialsumme (Teilsumme). Die Partialsummen $S_0, S_1, S_2, \dots$ bilden wieder eine Folge reeller Zahlen. Konvergiert diese gegen einen Grenzwert S ,

d.h.,

$$S_n \rightarrow S, \quad \text{für } n \rightarrow \infty,$$

so sagt man, dass die unendliche Reihe konvergiert. Den Grenzwert bezeichnet man dann mit

$$\sum_{k=0}^{\infty} a_k = S.$$

Konvergiert die Folge der Teilsummen nicht, so hat die unendliche Reihe keinen Wert und heißt divergent.

► 1.7.4 Die (endliche) geometrische Reihe

▷ Summenformel der geometrischen Reihe

Für alle $x \in \mathbb{R} \setminus \{1\}$ und $n \in \mathbb{N}$ gilt:

$$1 + x + x^2 + \cdots + x^n = \sum_{i=1}^n x^i = \frac{1 - x^{n+1}}{1 - x}.$$

Herleitung: Es gilt:

$$\begin{aligned} (1 + x + x^2 + \cdots + x^n)(1 - x) &= 1 + x + x^2 + \cdots + x^n \\ &\quad - x - x^2 - x^3 - \cdots - x^{n+1} \\ &= 1 - x^{n+1}. \end{aligned}$$

Da $x \neq 1$, können wir beide Seiten durch $1 - x$ dividieren:

$$1 + x + \cdots + x^n = \frac{1 - x^{n+1}}{1 - x}$$

▷ Grenzwert der geometrischen Reihe

Für alle $x \in \mathbb{R}$ mit $|x| < 1$ gilt:

$$\lim_{n \rightarrow \infty} \sum_{i=1}^n x^i = \frac{1}{1 - x}.$$

Herleitung: Für $|x| < 1$ ist x^{n+1} eine Nullfolge. Daher folgt:

$$\sum_{i=0}^n x^i = \frac{1 - x^{n+1}}{1 - x} \rightarrow \frac{1 - 0}{1 - x} = \frac{1}{1 - x}.$$

Fortsetzung der Motivation:

Kommen wir zur Menge V_n der Substanz in der Zelle zurück:

$$V_n = v_0 \sum_{i=0}^n q^i, \quad \text{mit } q = 1 - p.$$

Da die geometrische Summe konvergent ist (bei uns ist $0 < q < 1$), konvergiert V_n :

$$V_n \rightarrow \frac{v_0}{1 - q},$$

wenn $n \rightarrow \infty$.

Zahlenbeispiel: Die Ausgangsmenge betrage $v_0 = 2$ [ml]. Pro Zeiteinheit werden 20% verbraucht. Dann nähert sich die Menge der Substanz in der Zelle im Zeitablauf dem Wert

$$\frac{v_0}{1 - q} = \frac{2}{0.8} = 2.5$$

an.

1.8**1.8 Funktionen und Abbildungen**

Wir haben schon an einigen Stellen mit Funktionen zu tun gehabt, ohne eine formale Definition angegeben zu haben. Um die wichtigen Begriffe *Umkehrfunktion*, *Stetigkeit* und *Differenzierbarkeit* einführen zu können, müssen wir das nachholen:

Sei $D \subset \mathbb{R}$. f heißt **Funktion** von D nach $\mathbb{R}$, i.Z.: $f : D \rightarrow \mathbb{R}$, wenn jedem **Argument** $x \in D$ genau ein Bildelement (Bild, Funktionswert)

$$y = f(x) \in \mathbb{R}$$

zugeordnet wird. D heißt **Definitionsbereich** und

$$W = f(D) = \{f(x) : x \in D\}$$

Wertebereich von f . Die Menge aller Paare $(x, f(x))$ für $x \in D$ nennt man den **Graph** von f .

Eine erste wichtige Eigenschaft einer Funktion ist ihr Monotonieverhalten. Eine Funktion f heißt **monoton wachsend**, wenn aus $x_1 \leq x_2$ mit $x_1, x_2 \in D$ folgt: $f(x_1) \leq f(x_2)$, also wenn die Ungleichheitsrelation $\leq$ von f respektiert wird:

$$x_1 \leq x_2 \Rightarrow f(x_1) \leq f(x_2).$$

f heißt **streng monoton wachsend**, wenn gilt: $x_1 < x_2 \Rightarrow f(x_1) < f(x_2)$.
 f heißt **monoton fallend**, wenn für alle $x_1, x_2 \in D$ mit $x_1 \geq x_2$ folgt:
 $f(x_1) \geq f(x_2)$. f heißt **streng monoton fallend**, wenn für alle $x_1, x_2 \in D$
mit $x_1 < x_2$ folgt: $f(x_1) > f(x_2)$.

Nicht immer interessieren nur Zuordnungen zwischen Mengen von Zahlen. Allgemeiner nennt man eine Zuordnung T , die jedem Element einer Menge A auf eindeutige Weise ein Element aus irgendeiner anderen Menge B zuordnet eine **Abbildung**, i.Z.: $T : A \rightarrow B$. A heißt **Urbildmenge**, B **Bildmenge**. Funktionen sind also Abbildungen.

Ein interessantes und wichtiges Beispiel einer Abbildung ist der genetische Code.

Beispiel 1.8.1 Der genetische Code:
1.8.1

Die DNA ist ein Doppelstrang (Doppelhelix), der aus vier verschiedenen Desoxyribonukleotiden¹ zusammengesetzt ist, die sich dadurch unterscheiden, dass sie vier verschiedene Basen enthalten: Adenin (A), Thymin (T), Guanin (G) und Cytosin (C). (Die RNA unterscheidet sich von der DNA durch einen anderen Zucker. Zudem tritt statt Thymin Uracil (U) auf). Der genetische Code wird durch 64 Nukleotid-Tripel realisiert, von denen 61 für 20 Aminosäuren kodieren. Ein solches Tripel nennt man auch **Kodon**. Hierdurch wird eine Abbildung

$$G : K \rightarrow A$$

von der Menge

$$K = \{abc : a, b, c \in \{U, C, A, G\}\} = \{UUU, UUG, \dots, GGG\}$$

der 64 dreistelligen Nukleotid-Sequenzen in die Menge der 20 natürlich vorkommenden Aminosäuren

$$A = \{\text{Ala}, \text{Arg}, \dots, \text{Val}\}$$

gegeben (s. Tabelle 1.1 und Tabelle 1.2). Für eine Sequenz $S_n = a_1 a_2 a_3 \dots a_n$ beliebiger Länge definiert man einfach

$$G(S_n) = G(a_1 a_2 a_3) G(a_4 a_5 a_6) \dots G(a_{n-2} a_{n-1} a_n).$$

Dann ist bspw.

$$G(UCUCAGUCU) = \text{Ser Gln Ser}$$

¹Desoxyadenosinphosphat, Desoxythyminphosphat, Desoxyguanosinphosphat und Desoxycytosinphosphat.

2nd					
	U	C	A	G	
1st					3rd
U	Phe	Ser	Tyr	Cys	U
	Phe	Ser	Tyr	Cys	C
	Leu	Ser	TC	TC	A
	Leu	Ser	TC	Trp	G
C	Leu	Pro	His	Arg	U
	Leu	Pro	His	Arg	C
	Leu	Pro	Gln	Arg	A
	Leu	Pro	Gln	Arg	G
A	Ile	Thr	Asn	Ser	U
	Ile	Thr	Asn	Ser	C
	Ile	Thr	Lys	Arg	A
	Met	Thr	Lys	Arg	G
G	Val	Ala	Asp	Gly	U
	Val	Ala	Asp	Gly	C
	Val	Ala	Glu	Gly	A
	Val	Ala	Glu	Gly	G

Tabelle 1.1: Der genetische Code: 61 Nukleotid-Tripel (Kodons) kodieren für 20 Aminosäuren. TC bezeichnet Stopp-Kodons.

Aminosäure (engl.)	Code 1	Code 2
alanine	Ala	A
arginine	Arg	R
aspartic acid	Asp	D
asparagine	Asn	N
cysteine	Cys	C
glutamic acid	Glu	E
glutamine	Gln	Q
glycine	Gly	G
histidine	His	H
isoleucine	Ile	I
leucine	Leu	L
lysine	Lys	K
methionine	Met	M
phenylalanine	Phe	F
proline	Pro	P
serine	Ser	S
threonine	Thr	T
tryptophan	Trp	W
tyrosine	Tyr	Y
valine	Val	V

Tabelle 1.2. Die 20 Aminosäuren und ihre englischen 3- bzw. 1-Buchstaben-Codierungen.

Ist eine Funktion $f : D \rightarrow E$ mit $E \subset \mathbb{R}$ vorgegeben, so stellen sich zwei Fragen:

1. Welche Elemente aus E werden von f angenommen?
2. Sind die Bildelemente von zwei verschiedenen Argumenten auch verschieden?

Diese Fragen geben Anlass zu zwei Definitionen: Eine Funktion f heißt **surjektiv**, wenn es zu jedem $y \in E$ ein $x \in D$ gibt mit $y = f(x)$, also wenn jedes Element aus E von f angenommen wird. Eine Funktion f heißt **injektiv**, wenn für alle $x_1, x_2 \in D$ mit $x_1 \neq x_2$ gilt: $f(x_1) \neq f(x_2)$. Bei einer injektiven Funktion sind also die Bilder verschiedener Argumente verschieden.

Der genetische Code ist surjektiv, aber nicht injektiv (warum?).

Ist eine Funktion $f : D \rightarrow E$ sowohl injektiv als auch surjektiv, so tritt jedes Element aus E als Bild auf und unterschiedlichen Argumenten aus D entsprechen unterschiedliche Bilder aus E . Die Zuordnung f heißt dann **bijektiv**. Ist f bijektiv, so ist jedem $x \in D$ genau ein (ein und nur ein) Element aus E zugeordnet, und umgekehrt.

Paradebeispiele für Bijektionen sind streng monoton wachsende oder streng monoton fallende Funktionen.

1.8.1 Anmerkung 1.8.1 Hier noch eine anschauliche, wenn auch blutrünstige Erklärung. Ein Indianer hat 10 Pfeile im Köcher, die er auf eine angreifende Truppe von Soldaten schießt. Er trifft jedesmal und jeder seiner Pfeile ist tödlich. Hierdurch wird eine Funktion f definiert, die jedem Pfeil die Nummer des getroffenen Soldaten zuordnet. f ist injektiv, wenn jeder Pfeil einen anderen Soldaten getroffen hat. f ist surjektiv, wenn alle Soldaten tot sind.

► 1.8.1 Komposition von Funktionen

Mitunter hat man es mit recht komplizierten Funktionen und Formeln zu tun, die nicht auf den ersten Blick zu verstehen sind. Das Generalrezept ist, sie in einzelne Bestandteile zu zerlegen. Die Funktionsvorschrift

$$h(x) = \sqrt{x^2 + 1}$$

besagt etwa, dass man zunächst $y = f(x) = x^2 + 1$ berechnet und anschließend $g(y) = \sqrt{y}$. Der Wertebereich von $f(x)$ ist $f(\mathbb{R}) = [1, \infty)$. Da dies eine Teilmenge des Definitionsbereichs von $g(y)$ ist, ist $h(x)$ für alle $x \in \mathbb{R}$ definiert. Also kann man schreiben

$$h(x) = g(f(x)), \quad x \in \mathbb{R}.$$

Ist allgemein f eine Funktion mit Definitionsbereich D und ist g eine Funktion, dessen Definitionsbereich das Bild $f(D)$ von f umfasst, so kann man für jedes $x \in D$ die Funktion g auf $f(x)$ anwenden, also die **Komposition** $g \circ f$,

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

bilden. Zunächst wird also die Funktion f auf das Argument x angewendet und man erhält $y = f(x)$. Auf y wird nun die Funktion g angewendet, was $z = g(y) = g(f(x))$ ergibt.

► 1.8.2 Umkehrfunktion

Mitunter sind funktionale Zusammenhänge zwischen zwei interessierenden (biologischen) Größen „falsch“ herum gegeben: Man kennt $y = f(x)$, weiß also, wie man für ein gegebenes x den zugehörigen y -Wert berechnet, hätte aber gern zu einem gegebenem y den zugehörigen x -Wert. Man möchte also die Funktion $y = f(x)$ *umkehren* zu einer Funktionsvorschrift $x = f^{-1}(y)$. Graphisch geschieht dies, indem man die Funktion f „anders herum“ abliest. Doch mitunter benötigt man eine explizite Formel.

Ist $f : D \rightarrow E$ eine bijektive Funktion, so existiert eine **Umkehrfunktion** $f^{-1} : f(D) \rightarrow D$, so dass

$$\begin{aligned} f^{-1}(f(x)) &= x && \text{für alle } x \in D \\ f(f^{-1}(y)) &= y && \text{für alle } y \in f(D) \end{aligned}$$

Graphisch ermittelt man die Umkehrfunktion durch Spiegelung an der Winkelhalbierenden. Rechnerisch erhält man f^{-1} durch Auflösen der Gleichung $y = f(x)$ nach x .

Schema:

$$\begin{aligned} y = f(x) &\Leftrightarrow \dots \\ &\Leftrightarrow x = \underbrace{\dots}_{=f^{-1}(y)} \end{aligned}$$

Ist eine Funktion in einem Intervall $[a, b]$ streng monoton, so existiert die Umkehrfunktion f^{-1} .

Beispiel 1.8.2 Die Funktion $f : [0, \infty) \rightarrow \mathbb{R}$, $y = x^2 + 4$, ist auf $D = [0, \infty)$ streng monoton wachsend. Das Bild von f ist $f(D) = [4, \infty)$. Es gilt für alle

1.8.2

$x \geq 0$:

$$\begin{aligned} y = x^2 + 4 &\Leftrightarrow y - 4 = x^2 \\ &\Leftrightarrow x = \sqrt{y - 4} \end{aligned}$$

Also ist $f^{-1}(y) = \sqrt{y - 4}$ mit Definitionsbereich $[4, \infty) = f(D)$.

1.9 Stetigkeit

► 1.9.1 Motivation

Enzymatische Reaktion: Enzyme sind an fast allen Stoffwechsel - Reaktionen beteiligt. So hängt beispielsweise bei enzymatischen Reaktionen die Aktivität y eines Enzyms von der Temperatur x ab, so dass wir $y = f(x)$ schreiben können, wenn alle anderen Einflussgrößen konstant gehalten werden. Grundsätzlich gilt, dass die Enzym-Aktivität mit steigender Temperatur zunimmt. Allerdings werden ab ca 50° Celsius die Enzyme zerstört. Somit hat die Funktion $f(x)$ eine Sprungstelle bei $x = 50$. Abbildung 1.9.1 zeigt eine hypothetische unstetige Aktivitätsfunktion. Die Enzym-Aktivität hängt auch

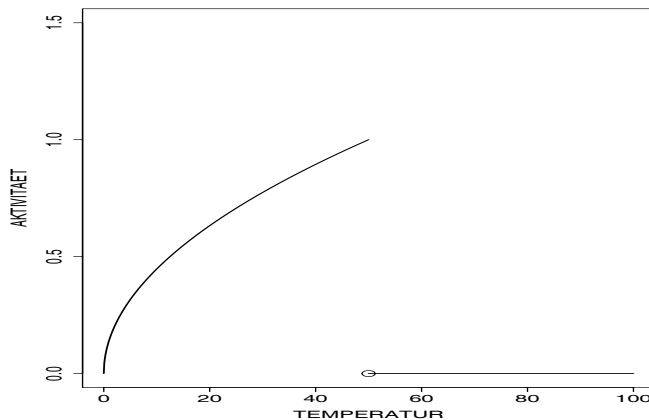


Abbildung 1.3. Eine hypothetische unstetige Aktivitätsfunktion

von anderen Größen ab, etwa dem pH-Wert: Amylase wirkt im Mund (pH-Wert: ca. 7). Gelangt a-Amylase jedoch in den Magen, so verliert sie wegen der geänderten pH-Bedingungen sofort ihre Aktivität.

➤ 1.9.2 Begriffsbildung

Es stellt sich die Frage, ob solche funktionalen Zusammenhänge stetig verlaufen oder sich auch abrupte Änderungen ergeben können. Diesen Sachverhalt kann man wie folgt präzisieren: Von Stetigkeit wollen wir sprechen, wenn die y -Änderung beliebig klein wird, wenn man die Variation des x -Werte immer kleiner wählt. Mathematisch ausgedrückt: Aus $x_n \rightarrow x_0$ für $n \rightarrow \infty$ (Konvergenz einer Folge von x -Werte gegen einen fest gewählten Wert x_0) soll $f(x_n) \rightarrow f(x_0)$ für $n \rightarrow \infty$ folgen (Konvergenz der zugehörigen y -Werte). Hier die genaue Definition:

Eine Funktion $f : D \rightarrow \mathbb{R}$ heißt **stetig im Punkt** x_0 , wenn für alle Folgen (x_n) mit $x_n \rightarrow x_0$, $n \rightarrow \infty$, gilt:

$$f(x_n) \rightarrow f(x_0), \quad n \rightarrow \infty.$$

f heißt stetig, wenn f stetig in allen Punkten $x \in D$ ist. Eine andere Schreibweise hierfür ist:

$$f(x) \rightarrow f(x_0) \quad \text{für } x \rightarrow x_0.$$

Man sagt: $f(x)$ konvergiert gegen $f(x_0)$, wenn x gegen x_0 konvergiert.

Bei dieser Definition sind zwei Dinge wichtig: $f(x_n)$ muss für *alle* Folgen konvergieren (ohne Ausnahme) und das Grenzelement muss mit $f(x_0)$ übereinstimmen.

Allgemein sagt man, dass $f(x)$ gegen $a \in \mathbb{R}$ konvergiert, wenn für alle Folgen (x_n) mit $x_n \rightarrow x_0$ für $n \rightarrow \infty$ folgt: $f(x_n) \rightarrow a$. Stetigkeit liegt vor, wenn zudem $a = f(x_0)$ gilt.

Beispiel 1.9.1 $f(x)$ sei die Anzahl der die Ziellinie passierenden Skifahrer zur Zeit x . Da Skifahrer in der Realität nie exakt zur selben Zeit im Ziel ankommen, ist $f(x)$ genau dann 1, wenn ein Skifahrer im Ziel eintrifft, und sonst 0. Ist x_0 solch ein Ankunftszeitpunkt, dann gilt für die Folge $x_n = x_0 + 1/n$, $n \in \mathbb{N}$: $f(x_n) = 0$ für alle n . Also folgt für diese spezielle Folge: $f(x_n) \rightarrow 0$, wenn $n \rightarrow \infty$. Da aber $f(x_0) = 1 \neq 0$, ist f nicht stetig in $x = x_0$.

1.9.1

Wir wollen für zwei Funktionen explizit zeigen, dass sie stetig sind.

Beispiel 1.9.2 Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$, $y = x^2$, ist stetig. Um dies einzusehen, sei $x_0 \in \mathbb{R}$ ein beliebiger Punkt aus $\mathbb{R}$ und $(x_n) \subset \mathbb{R}$ eine konvergente Folge mit Limes x_0 . Dann gilt nach den Rechenregeln für konvergente Reihen.

1.9.2

$$f(x_n) = x_n^2 = x_n \cdot x_n \rightarrow x_0 \cdot x_0 = x_0^2,$$

wenn $n \rightarrow \infty$.

Hier noch ein biologisches Beispiel:

1.9.3 Beispiel 1.9.3 Biologische Zusammenhänge zwischen einer Dosis x und der zugehörigen Wirkung y können oftmals gut durch die **Michaelis - Menten - Funktion** $f : [0, \infty) \rightarrow \mathbb{R}^+$,

$$y = f(x) = \frac{bx}{a+x},$$

beschrieben werden, wobei a und b zwei positive Konstanten sind. f heißt dann auch *Dosis-Wirkung-Funktion*. Diese Funktion tritt bei enzymatischen Reaktionen auf und beschreibt dort die Geschwindigkeit y der Reaktion in Abhängigkeit von der Substrat - Konzentration x . Wir wollen nachweisen, dass f eine stetige Funktion ist. Hierzu sei (x_n) eine Folge mit $x_n \rightarrow x_0 \neq -a$, wenn $n \rightarrow \infty$. Dann folgt

$$b \cdot x_n \rightarrow bx_0, \quad \text{für } n \rightarrow \infty$$

und

$$a + x_n \rightarrow a + x_0, \quad \text{für } n \rightarrow \infty.$$

Da $x_0 \neq -a$, ist ab einem Index n_0 stets $a + x_n > 0$ erfüllt. Dann folgt

$$\frac{b \cdot x_n}{a + x_n} \rightarrow \frac{b \cdot x_0}{a + x_0}, \quad \text{wenn } n \rightarrow \infty.$$

Die linke Seite ist gerade $f(x_n)$, die rechte Seite $f(x_0)$. Also gilt

$$f(x_0) = \lim_{n \rightarrow \infty} f(x_n)$$

für alle Folgen (x_n) mit $x_0 = \lim_{n \rightarrow \infty} x_n$. Damit ist die Stetigkeit von $f(x)$ gezeigt.

► 1.9.3 Eigenschaften stetiger Funktionen

Wir haben gesehen, dass die Dosis-Wirkung-Funktion aus Beispiel 1.9.3 stetig ist. Hier drei naheliegende Fragen:

1. Gibt es zu jeder Wirkung y eine Dosis x , die zu dieser Wirkung führt?
2. Kann die Dosis-Wirkung-Funktion so aufgeschrieben werden, dass man zu jeder Wirkung y die einzusetzende Dosis x erhält?
3. Nimmt die Dosis-Wirkung-Funktion für einen Dosierungsbereich $[a, b]$ ihr Minimum und Maximum an? D.h.: Gibt es Dosierungen $x_{\min}$ und $x_{\max}$

zwischen a und b , so dass $f(x_{\min})$ genau die minimale und $f(x_{\max})$ die maximale Wirkung über diesen Dosierungsbereich ist?

Zeichnet man die Dosis-Wirkungs-Funktion, so suggeriert der Graph, dass alle drei Fragen positiv zu beantworten sind. Dies liegt jedoch nicht an der speziellen Form der Dosis-Wirkungs-Funktion, sondern an ihren *qualitativen Eigenschaften*: Stetigkeit und Monotonie. Grundlage dieser Erkenntnis sind die folgenden wichtigen Eigenschaften stetiger Funktionen.

1. Zwischenwertsatz: Ist $f : [a, b] \rightarrow \mathbb{R}$ stetig, so gibt es zu jedem y mit $f(a) \leq y \leq f(b)$ ein $x \in [a, b]$ mit $f(x) = y$.
2. $f : D \rightarrow \mathbb{R}$ sei stetig und streng monoton. Dann existiert eine stetige und streng monotone Umkehrfunktion $f^{-1} : f(D) \rightarrow D$ mit

$$f^{-1}(f(x)) = x, \quad x \in D$$

und

$$f(f^{-1}(y)) = y, \quad y \in f(D).$$

3. Jede in einem abgeschlossenen Intervall $[a, b]$ stetige Funktion ist dort beschränkt und nimmt ihr Maximum und Minimum an.

1.10 Exponentialfunktion

1.10

Die Exponentialfunktion ist von fundamentaler Bedeutung. Sie verallgemeinert die Potenzbildung a^x (als Funktion von x) auf beliebige reelle Exponenten x . Potenzen a^x mit ganzen Exponenten hatten eine entscheidende Rolle bei *zeit-diskreten* Wachstumsprozessen gespielt. Die Exponentialfunktion tritt nun bei *kontinuierlichen* Wachstumsprozessen auf. Zeit-diskret heißt, dass die relevanten Zeitpunkte einzelne, isolierte Zeitpunkte sind, etwa $1, 2, 3, 4, \dots$. Kontinuierlich (zeit-stetig) meint, dass die Zeit ein Intervall der Form $[a, b]$ (oder auch $[a, \infty)$) durchläuft.

➤ 1.10.1 Definition

Potenzen der Form $f(x) = a^x$ für eine beliebige Basis $a > 0$ und rationale Exponenten $x \in \mathbb{Q}$ waren in drei Schritten definiert worden.

1. für $x \in \mathbb{N}$: $a^x = a \cdot \dots \cdot a$ (n -mal).

Erweiterung auf negative Exponenten:

$$2. \text{ für } x \in \mathbb{N}: \quad a^{-x} = \frac{1}{a^x}, \quad a^0 = 1.$$

Und schließlich

$$3. \text{ für } x = p/q \in \mathbb{Q}: \quad a^{p/q} = \sqrt[q]{a^p}.$$

Was wir brauchen, ist die Erweiterung auf beliebige reelle Exponenten $x \in \mathbb{R}$. Wir können jede reelle Zahl x durch eine Folge von Brüchen (x_n) annähern, etwa indem wir in der Dezimalbruch-Darstellung nach der n -ten Stelle abbrechen - z.B.:

$$x = 1.1415\dots, \quad x_0 = 1, x_1 = 1.1, x_2 = 1.14, x_3 = 1.141, \text{ etc.}$$

Formal: Sei (x_n) eine Folge rationaler Zahlen mit

$$x = \lim_{n \rightarrow \infty} x_n.$$

Für jedes Element der annähernden Folge können wir die Potenz a^{x_n} nach obigen Regeln berechnen. Es ist nun nahe liegend, $f(x)$ als Grenzwert der Folge der Bilder

$$f(x_0) = a^{x_0}, f(x_1) = a^{x_1}, \dots, f(x_n) = a^{x_n}, \dots$$

zu definieren, sofern dieser existiert. Da dieser Grenzübergang gültig ist, kann man in der Tat die Festsetzung

$$f(x) = a^x := \lim_{n \rightarrow \infty} a^{x_n}$$

treffen. Diese Funktion heißt **Exponentialfunktion zur Basis a** :

$$\exp_a : \mathbb{R} \rightarrow \mathbb{R}^+, \quad \exp_a(x) = a^x.$$

► 1.10.2 Eigenschaften

Die schon formulierten Rechenregeln für Potenzen übertragen sich auf die Funktion $\exp_a(x)$.

1. **Fundamentalgleichung:** Für alle $x, y \in \mathbb{R}$ und $a > 0$ gilt:

$$\exp_a(x + y) = \exp_a(x) \cdot \exp_a(y)$$

$$2. \exp_a(x) \cdot \exp_b(x) = \exp_{ab}(x).$$

$$3. \exp_a(x)^y = \exp_a(xy)$$

$$4. \exp_a(x) \text{ ist streng monoton wachsend.}$$

1.11 Kontinuierliches Wachstum

1.11

Bei vielen Wachstumsvorgängen ist es realistischer von einer zeit-stetigen Entwicklung auszugehen, anstatt von einer zeit-diskreten.

Um zeit-stetige Wachstumsprozesse unter konstanten Wachstumsbedingungen aus zeit-diskreten Überlegungen abzuleiten, wollen wir annehmen, dass für kleine Zeitabstände $\Delta t > 0$ das Populationswachstum näherungsweise proportional zur Größe der Population und zur Zeitspanne Δt ist. D.h.

$$y(t + \Delta t) \approx y(t) + \lambda y(t) \Delta t = (1 + \lambda \Delta t) y(t),$$

oder - äquivalent - dass die Wachstumsrate proportional zu Zeitspanne aber zeitlich konstant ist:

$$\frac{y(t + \Delta t) - y(t)}{\Delta t} \approx \lambda y(t).$$

Die Proportionalitätskonstante λ heißt auch **Intensität**. Für einen festen Zeitpunkt $t > 0$ zerlegen wir nun das Zeitintervall $[0, t]$ in n gleichlange Intervalle $[t_0, t_1], [t_1, t_2], \dots, [t_{n-1}, t_n]$ der Länge $\Delta t = t/n$. D.h.:

$$t_k = k \cdot \Delta t, k = 1, \dots, n, \quad t_0 = 0, t_n = t.$$

Dann ist $t_n = t_{n-1} + \Delta t$, $(1 + \lambda \Delta t)^n = (1 + \frac{\lambda t}{n})^n$. Also folgt:

$$\begin{aligned} y(t) &= y(t_n) \\ &\approx y(t_{n-1}) \left(1 + \frac{\lambda t}{n}\right) \\ &\approx y(t_{n-2}) \left(1 + \frac{\lambda t}{n}\right)^2 \\ &\vdots \\ &\approx y(t_0) \left(1 + \frac{\lambda t}{n}\right)^n \end{aligned}$$

Was passiert nun, wenn wir die Anzahl der Teilintervalle n gegen unendlich streben lassen? Dann geht die Näherung $y(t + \Delta t) \approx (1 + \lambda \Delta t) y(t)$ wegen $\Delta t = t/n \rightarrow 0$ in eine Gleichheit über. Also:

$$y(t) = y(t_0) \cdot \lim_{n \rightarrow \infty} \left(1 + \frac{\lambda t}{n}\right)^n.$$

Es gilt nun

$$e^x = \lim_{n \rightarrow \infty} \left(1 + \frac{x}{n}\right)^n.$$

Für $x = 1$ erhält man die Euler'sche Zahl $e = 2.71828 \dots$

Wir erhalten also als Ergebnis:

$$y(t) = y(t_0) \cdot e^{\lambda t}.$$

Hierbei ist $e^{\lambda t}$ der zeit-stetige Wachstumsfaktor und λ die zeitlich konstante Intensität.

Betrachtet man nicht das Zeitintervall $[0, t]$, sondern etwas praxisnäher das Intervall $[t_0, t]$, so schreibt sich das exponentielle Wachstumsgesetz in der Form

$$y(t) = y(t_0)e^{\lambda(t-t_0)}.$$

Das radioaktive Zerfallsgesetz

Der Zerfall radioaktiver Substanzen erfolgt in sehr guter Näherung nach einem exponentiellen Gesetz, das man hier üblicherweise in der Form

$$y(t) = y(t_0)e^{-\lambda(t-t_0)}$$

aufschreibt. Der Parameter λ heißt **Zerfallskonstante**. Üblicherweise gibt man jedoch nicht die Zerfallskonstante, sondern die **Halbwertszeit** T_H . Die Halbwertszeit ist diejenige Zeitspanne, nach der die Hälfte des Materials verstrahlt ist. Also gilt

$$\frac{y(t_0 + T_H)}{y(t_0)} = \frac{1}{2} \Leftrightarrow T_H = \frac{\ln 2}{\lambda}.$$

Hierbei ist $\ln 2$ diejenige reelle Zahl x mit $e^x = 2$ (s.u.).

1.11.1 Beispiel 1.11.1 Den Parameter λ kann man wie folgt aus Laborwerten in grober Näherung so bestimmen: Ein Zellhaufen habe sich während einer Zeiteinheit von $x(0) = 10$ auf $x(1) = 12$ Mengeneinheiten vermehrt. Es ist

$$x(1) - x(0) \approx \lambda \cdot x(0).$$

D.h.: $x(1) \approx x(0)(1 + \lambda)$. Folglich: $\lambda \approx \frac{x(1)}{x(0)} - 1 = 1.2 - 1 = 0.2$.

1.12 Der Logarithmus

Die Umkehrfunktion der Exponentialfunktion $\exp_a : \mathbb{R} \rightarrow (0, \infty)$ heißt **Logarithmus zur Basis a** und wird mit

$$\log_a : (0, \infty) \rightarrow \mathbb{R}$$

bezeichnet. Der Logarithmus zur Basis e heißt **natürlicher Logarithmus** und wird mit

$$\ln(x) = \log_e(x)$$

bezeichnet. Es gilt

$$y = \exp_a(x) = a^x \Leftrightarrow \log_a(y) = x.$$

Merkregel: Der Logarithmus zur Basis a extrahiert aus einem Potenzausdruck a^x den Exponenten x . Daher gilt auch

$$\log_a(1) = \log_a(a^0) = 0 \quad \text{und} \quad \log_a(a) = \log_a(a^1) = 1.$$

Ferner kann man jede reelle Zahl schreiben als:

$$y = \exp_a(\log_a(y)) = a^{\log_a(y)}$$

➤ 1.12.1 Rechenregeln

1. $\log(xy) = \log(x) + \log(y)$
2. $\log(x/y) = \log(x) - \log(y)$
3. $\log(x^y) = y \log(x)$
4. Umrechnen von Logarithmen:

$$\log_a(x) = \log_a(b) \cdot \log_b(x)$$

Beispiel 1.12.1 Das Modell des gleichmäßigen konstanten Wachstums lautet:

1.12.1

$$x(t) = x(t_0)e^{\lambda t}.$$

Wir wollen nun diese Gleichung nach λ auflösen:

$$\begin{aligned} x(t) &= x(t_0)e^{\lambda t} \\ \Leftrightarrow \frac{x(t)}{x(t_0)} &= e^{\lambda t} \\ \Leftrightarrow \lambda t &= \ln \frac{x(t)}{x(t_0)} \\ \Leftrightarrow \lambda &= \frac{1}{t} \ln \frac{x(t)}{x(t_0)} \end{aligned}$$

Für $x(1) = 12$ und $x(0) = 10$ erhalten wir $\lambda = \ln 1.2 = 0.18$.

Mathematische Grundlagen der empirischen Forschung

Steland, A.

2004, XII, 376 S. 28 Abb., Softcover

ISBN: 978-3-540-03700-2