

4 Erweiterungen der natürlichen Zahlen

1. Ganze Zahlen

Die arithmetischen Operationen der Addition und Multiplikation sind in den natürlichen Zahlen nur eingeschränkt umkehrbar. Will man zu jedem n ein additives Inverses, also eine Zahl $(-n)$ mit $n + (-n) = 0$, so muss man negative Zahlen einführen. Man kann nun diese Erweiterung von $\mathbb{N}$ aus $\mathbb{N}$ heraus konstruieren: Definiere auf dem kartesischen Produkt $\mathbb{N} \times \mathbb{N}$ eine Addition durch komponentenweises Addieren,

$$(m_1, m_2) + (n_1, n_2) := (m_1 + n_1, m_2 + n_2)$$

und betrachte die Äquivalenzrelation

$$(m_1, m_2) \sim (n_1, n_2) \Leftrightarrow m_1 + n_2 = m_2 + n_1.$$

Symmetrie und Reflexivität sind klar, aber auch Transitivität ist leicht nachzuprüfen:

$$(m_1, m_2) \sim (n_1, n_2) \Leftrightarrow m_1 + n_2 = m_2 + n_1$$

$$(n_1, n_2) \sim (k_1, k_2) \Leftrightarrow n_1 + k_2 = n_2 + k_1$$

Addiert man die Gleichungen, erhält man

$$m_1 + k_2 + n_1 + n_2 = m_2 + k_1 + n_1 + n_2$$

oder äquivalent $m_1 + k_2 = m_2 + k_1$ und damit $(m_1, m_2) \sim (k_1, k_2)$. Damit zerfällt $\mathbb{N} \times \mathbb{N}$ in disjunkte Äquivalenzklassen

$$[m_1, m_2] := \{(n_1, n_2) \mid m_1 + n_2 = m_2 + n_1\}.$$

Sind (m_1, m_2) und (n_1, n_2) zwei Paare, so ist die Äquivalenzklasse der Summe $(m_1 + n_1, m_2 + n_2)$ schon durch die Äquivalenzklassen $[m_1, m_2]$ und $[n_1, n_2]$ festgelegt: Sei etwa $(k_1, k_2) \sim (n_1, n_2)$, also $n_1 + k_2 = n_2 + k_1$. Dann gilt

$$(m_1 + n_1) + (m_2 + k_2) = m_1 + m_2 + n_1 + k_2 = m_1 + m_2 + n_2 + k_1 = (m_2 + n_2) + (m_1 + k_1),$$

also $(m_1 + n_1, m_2 + n_2) \sim (m_1 + k_1, m_2 + k_2)$. Folglich kann man auf der Menge der Äquivalenzklassen durch

$$[m_1, m_2] \oplus [n_1, n_2] := [m_1 + n_1, m_2 + n_2]$$

eine „Addition“ erklären, die kommutativ und assoziativ ist. Durch wiederholtes Anwenden der Vorgängerfunktion kann man für jede Äquivalenzklasse einen Repräsentanten finden, bei dem eine der beiden Komponenten 0 ist: Eine Klasse $[n_1, n_2]$ mit $n_1 > n_2$ hat einen Repräsentanten der Form $(n, 0)$ mit $n > 0$, und ist $n_1 < n_2$, so gibt es einen Repräsentanten der Form $(0, n)$ mit $n > 0$. Die Paare der Form $[n, 0]$ repräsentieren die natürlichen Zahlen, und die Paare $[0, m]$ mit $m > 0$ die negativen Zahlen.

Neutrales Element der Addition $\oplus$ ist die Klasse $[0, 0]$. Zu der Klasse $[n_1, n_2]$ sei

$$-[n_1, n_2] := [n_2, n_1].$$

Dann gilt $-[n_1, n_2] \oplus [n_1, n_2] = [0, 0]$, d.h. jede Klasse hat ein eindeutig bestimmtes Inverses für die Operation $\oplus$. Es ist üblich, die Operation „ $\oplus - [n_1, n_2]$ “ als *Subtraktion* der Klasse $[n_1, n_2]$ zu bezeichnen. Damit ist die Konstruktion der ganzen Zahlen aus den natürlichen Zahlen abgeschlossen:

$$\mathbb{Z} := (\mathbb{N} \times \mathbb{N}) / \sim \quad \text{mit der Addition } \oplus.$$

Um wieder zu der gewohnten Notation zu gelangen, schreibt man für $n \in \mathbb{N}$ die Klasse $[n, 0]$ wieder als n , die Klasse $[0, m]$ als $(-m)$ und verwendet das übliche Symbol $+$ für die Addition.

Die ganzen Zahlen zusammen mit der auf ihnen erklärten Addition sind ein Beispiel für ein ganz fundamentales Konzept der Mathematik.

Sei G eine Menge. Unter einer Verknüpfung (oder auch Operation) auf G versteht man eine Abbildung

$$*: G \times G \rightarrow G.$$

Statt $*(x, y)$ schreibt man auch $x*y$. Oft lässt man das Zeichen $*$ auch ganz weg und schreibt nur xy für die Verknüpfung von $x \in G$ und $y \in G$; dies ist die sogenannte **multiplikative Schreibweise**.

Definition 1.1. Eine **Gruppe** ist ein Tripel $(G, *, e)$ bestehend aus einer Menge G , einer Verknüpfung $*: G \times G \rightarrow G$ und einem Element $e \in G$, genannt **neutrales Element**, so dass gilt:

G1 Für alle $x, y, z \in G$ ist $(x * y) * z = x * (y * z)$ (**Assoziativität**).

G2 Für alle $x \in G$ ist $e * x = x * e = x$ (**neutrales Element**).

G3 Zu jedem $x \in G$ gibt es ein $x' \in G$ mit $x * x' = x' * x = e$ (**Existenz des Inversen**).

Gilt zusätzlich noch

G4 Für alle $x, y \in G$ ist $x * y = y * x$ (**Kommutativität**),

so heißt G eine **abelsche** oder **kommutative Gruppe**.

Beispiele. 1. Die Addition macht $\mathbb{Z}$ zu einer abelschen Gruppe mit neutralem Element 0.

2. Sei M eine (endliche) Menge, und sei

$$\Sigma(M) := \{f: M \rightarrow M \mid f \text{ ist bijektiv}\}$$

die Menge aller bijektiven Selbstabbildungen von M . Die Komposition „ $\circ$ “ von Abbildungen macht $(\Sigma(M), \circ, \text{id}_M)$ zu einer Gruppe:

G1 Sind $f, g, h: M \rightarrow M$ Abbildungen, so gilt $(f \circ g) \circ h = f \circ (g \circ h)$, denn die Komposition von Abbildungen ist assoziativ.

G2 Für alle Abbildungen $f: M \rightarrow M$ gilt $f \circ \text{id}_M = \text{id}_M \circ f = f$.

G3 Jede bijektive Abbildung $f: M \rightarrow M$ hat eine Inverse f' (siehe Lemma 2.5 in Kapitel 2).

Man nennt $\Sigma(M)$ auch die Gruppe der Permutationen von M . In Satz 3.4 von Kapitel 3 wurde gezeigt, dass es $n!$ Elemente in $\Sigma(M)$ gibt, wenn M aus n Elementen besteht. Die Gruppe $\Sigma(M)$ ist im allgemeinen nicht abelsch: Sei $M = \{1, 2, 3\}$ und seien $f, g: M \rightarrow M$ gegeben durch

$$f(1) = 2, f(2) = 1, f(3) = 3 \quad \text{bzw.} \quad g(1) = 1, g(2) = 3, g(3) = 2.$$

Dann gilt

$$\begin{aligned} (g \circ f)(1) &= 3, (g \circ f)(2) = 1, (g \circ f)(3) = 2, \\ \text{aber } (f \circ g)(1) &= 2, (f \circ g)(2) = 3, (f \circ g)(3) = 1. \end{aligned}$$

3. Die natürlichen Zahlen bilden mit der Addition keine Gruppe, denn das Axiom **G3** ist verletzt: es gibt keine Inversen.

Bei einer abelschen Gruppe $(G, *, e)$ schreibt man die Verknüpfung oft **additiv**, also $x + y$ statt $x * y$ und 0 für das neutrale Element e .

Es bleibt noch die Multiplikation auf $\mathbb{Z}$ zu erklären: Man setzt

$$[m_1, m_2] \odot [n_1, n_2] := [m_1 \cdot n_1 + m_2 \cdot n_2, m_1 \cdot n_2 + m_2 \cdot n_1].$$

(Diese Formel ergibt sich aus der Interpretation von $[m_1, m_2]$ als $m_1 - m_2$.) Es lässt sich wieder einfach nachrechnen, dass diese Multiplikation wohldefiniert ist, also die rechte Seite nicht von der Wahl der Repräsentanten auf der linken Seite abhängt. Ebenso prüft man nach, dass die Operation assoziativ und kommutativ ist, und dass die Distributivgesetze gelten. Das neutrale Element der Multiplikation ist die Klasse $1 = [1, 0]$. Nachdem dies geklärt ist, kann man wieder zur gewohnten Notation für die Multiplikation zurückkehren.

Auf der Menge $\mathbb{Z}$ gibt es also zwei Operationen, deren Vertauschungsregeln die Distributivgesetze sind. Eine solche Menge nennt man einen **Ring**:

Definition 1.2. Ein **Ring** ist ein Quintupel $(R, +, \cdot, 0, 1)$ bestehend aus einer Menge R , einer Verknüpfung $+: R \times R \rightarrow R$, genannt Addition, einer Verknüpfung $\cdot: R \times R \rightarrow R$, genannt Multiplikation, und zwei Elementen $0 \in R$ und $1 \in R$ mit $1 \neq 0$, so dass gilt:

R1 Das Tripel $(R, +, 0)$ ist eine abelsche Gruppe.

R2 Für alle $x, y, z \in R$ ist

$$x \cdot (y \cdot z) = (x \cdot y) \cdot z$$

(Assoziativität der Multiplikation).

R3 Für alle $x \in R$ gilt $x \cdot 1 = 1 \cdot x = x$ (Eigenschaft des Einselements)

R4 Für alle $x, y, z \in R$ gelten die Distributivgesetze

$$x \cdot (y + z) = x \cdot y + x \cdot z, \quad (x + y) \cdot z = x \cdot z + y \cdot z.$$

Gilt zusätzlich noch das Kommutativgesetz für die Multiplikation, also

R5 für alle $x, y \in R$ ist $x \cdot y = y \cdot x$,

so heißt R ein kommutativer Ring.

Insbesondere wird für die Multiplikation nicht verlangt, dass es Inverse gibt.

Beispiele. 1. Die ganzen Zahlen bilden mit (gewöhnlicher) Addition und Multiplikation einen kommutativen Ring.

2. Ein künstlich anmutendes Beispiel eines nichtkommutativen Ringes: Sei $R = \mathbb{Z}^4$ mit Addition

$$(x_1, x_2, x_3, x_4) + (y_1, y_2, y_3, y_4) := (x_1 + y_1, x_2 + y_2, x_3 + y_3, x_4 + y_4)$$

(man nennt dies die komponentenweise Addition) und Multiplikation

$$(x_1, x_2, x_3, x_4) \cdot (y_1, y_2, y_3, y_4) := (x_1 y_1 + x_2 y_3, x_1 y_2 + x_2 y_4, x_3 y_1 + x_4 y_3, x_3 y_2 + x_4 y_4).$$

Das Nullelement ist $(0, 0, 0, 0)$ und das Einselement $(1, 0, 0, 1)$. Es ist dann nicht besonders schwierig (wenn auch langwierig), die Ringaxiome nachzuprüfen. Aber diese Multiplikation ist nicht kommutativ: es gilt

$$(1, 0, 0, 0) \cdot (0, 1, 0, 0) = (0, 1, 0, 0) \\ \text{aber } (0, 1, 0, 0) \cdot (1, 0, 0, 0) = (0, 0, 0, 0).$$

(Dieses Beispiel ist aus der Linearen Algebra geklaut, wie die Leser bald feststellen werden.)

Schließlich definiert man noch eine Totalordnung durch

$$[m_1, m_2] \leq [n_1, n_2] \quad \Leftrightarrow \quad m_1 + n_2 \leq m_2 + n_1.$$

Diese Ordnung ist allerdings keine Wohlordnung mehr, denn es gibt Teilmengen von $\mathbb{Z}$ (wie etwa $\mathbb{Z}$ selbst), die kein kleinstes Element haben. Aber zumindest hat noch jede nach unten beschränkte Teilmenge ein kleinstes Element: ist n_0 eine untere Schranke von $M \subset \mathbb{Z}$, so betrachte die Menge $M' := \{n \in \mathbb{Z} \mid n + (-n_0) \in M\}$. Dies ist eine Teilmenge von $\mathbb{N}$ und besitzt folglich ein kleinstes Element.

In der hier präsentierten Konstruktion der ganzen Zahlen ist $\mathbb{N}$ keine Teilmenge von $\mathbb{Z}$, denn $\mathbb{Z}$ ist ja eine Menge von Äquivalenzklassen von Paaren. Aber trotzdem findet man $\mathbb{N}$ mit seinen arithmetischen Operationen wieder: Definiere die Abbildung $I: \mathbb{N} \rightarrow \mathbb{Z}$ durch

$$I(n) := [n, 0].$$

Dann ist I eine injektive Abbildung, die Addition, Multiplikation und die Ordnungsrelation erhält, d.h. so dass gilt:

$$\begin{aligned} I(m+n) &= I(m) \oplus I(n) \\ I(m \cdot n) &= I(m) \odot I(n) \\ m \leq n &\Leftrightarrow I(m) \leq I(n) \end{aligned}$$

Schreibt man $\mathbb{N}'$ für das Bild von I , so nennt man eine solche Abbildung einen Isomorphismus von $(\mathbb{N}, +, \cdot, \leq)$ auf $(\mathbb{N}', \oplus, \odot, \leq)$.

Der **Absolutbetrag** einer ganzen Zahl $n := [n_1, n_2]$ ist

$$|n| := \begin{cases} n & \text{wenn } n_1 \geq n_2, \\ -n & \text{wenn } n_1 < n_2. \end{cases}$$

2. Rationale Zahlen

Ein ganz analoges Verfahren wendet man an, um die ganzen Zahlen zu den rationalen Zahlen zu erweitern. Seien dazu jetzt $\mathbb{N}$ und $\mathbb{Z}$ (mit den üblichen Symbolen für Addition und Multiplikation) als gegeben angesehen, sowie $\mathbb{N}$ vermöge des Isomorphismus I mit der Teilmenge $\mathbb{N}' \subset \mathbb{Z}$ identifiziert.

Man betrachtet auf der Menge $\mathbb{Z} \times \mathbb{N}_+$ die Äquivalenzrelation

$$(m_1, m_2) \sim (n_1, n_2) \Leftrightarrow m_1 \cdot n_2 = m_2 \cdot n_1$$

und nennt die Äquivalenzklassen rationale Zahlen. Zu interpretieren ist dabei $[m_1, m_2]$ als der Bruch $\frac{m_1}{m_2}$. Damit ist auch vorgegeben, wie man die arithmetischen Operationen zu erklären hat:

$$\begin{aligned} [m_1, m_2] \oplus [n_1, n_2] &:= [m_1 \cdot n_2 + m_2 \cdot n_1, m_2 \cdot n_2] \\ [m_1, m_2] \odot [n_1, n_2] &:= [m_1 \cdot n_1, m_2 \cdot n_2] \end{aligned}$$

Dies sind nichts anderes als die wohlbekannten Regeln des Bruchrechnens. Die notwendigen Überprüfungen sind wieder elementar und langweilig (wenn auch langwierig) und werden deshalb ausgespart. Das neutrale Element der Addition ist die Klasse $[0, 1]$, das der Multiplikation die Klasse $[1, 1]$. Außerdem gibt es zu jedem von $[0, 1]$ verschiedenen Element $[p, q]$ ein multiplikatives Inverses, nämlich die Klasse

$$[p, q]^{-1} := \begin{cases} [q, p] & \text{falls } p > 0, \\ [-q, -p] & \text{falls } p < 0. \end{cases}$$

Schließlich noch die Ordnungsrelation: Man definiert zuerst eine Klasse $[p, q]$ als *positiv*, wenn es einen Repräsentanten mit $p > 0$ gibt, und dann $x < y$, falls $y - x$ positiv ist. Die so konstruierte Menge von Äquivalenzklassen mit ihren Additions- und Multiplikationsoperationen bezeichnet man mit $\mathbb{Q}$. Formal ist auch hier wieder $\mathbb{Z}$ keine Teilmenge, aber die injektive Abbildung

$$\mathbb{Z} \rightarrow \mathbb{Q}, \quad n \mapsto [n, 1]$$

identifiziert $\mathbb{Z}$ mit einer isomorphen Kopie in $\mathbb{Q}$. Nun wirft man die ganze Hilfsnotation wieder weg, schreibt

$$\frac{p}{q} = [p, q]$$

und benutzt die üblichen Additions- und Multiplikationssymbole.

Die Rechenregeln der rationalen Zahlen werden durch den Begriff des **Körpers** axiomatisiert:

Definition 2.1. *Ein Körper ist eine Menge K zusammen mit zwei Operationen*

$$\begin{aligned} +: K \times K &\rightarrow K, & (x, y) &\mapsto x + y & (\text{Addition}) \\ \cdot: K \times K &\rightarrow K, & (x, y) &\mapsto x \cdot y & (\text{Multiplikation}) \end{aligned}$$

und zwei Elementen $0 \in K$ und $1 \in K$, $1 \neq 0$, mit folgenden Eigenschaften: für alle $a, b, c \in K$ gilt

K1	$a + (b + c) = (a + b) + c$	$a \cdot (b \cdot c) = (a \cdot b) \cdot c$	(Assoziativität)
K2	$a + b = b + a$	$a \cdot b = b \cdot a$	(Kommutativität)
K3	$a \cdot (b + c) = a \cdot b + a \cdot c$		(Distributivität)
K4	$a + 0 = a$	$a \cdot 1 = a$	(neutrales Element)
K5	$\exists x \in K \ a + x = 0$	$a \neq 0 \Rightarrow \exists x \in K \ a \cdot x = 1$	(Inverse)

Insbesondere ist K also ein kommutativer Ring, hat aber zusätzlich die Eigenschaft, dass jedes Element $\neq 0$ ein multiplikatives Inverses besitzt. Mit anderen Worten: das Tripel $(K - \{0\}, \cdot, 1)$ ist ebenfalls eine abelsche Gruppe.

Nun gab es auf $\mathbb{Q}$ aber noch die Totalordnung $\leq$, mit den bekannten Rechenregeln für Ungleichungen. Diese Eigenschaften lassen sich ebenfalls als Axiome formulieren:

Definition 2.2. Ein Körper K heißt angeordnet, wenn es eine Totalordnung $\leq$ auf K gibt, so dass gilt:

A1 Sind $a > 0$ und $b > 0$, so folgt $a + b > 0$.

A2 Sind $a > 0$ und $b > 0$, so folgt $a \cdot b > 0$.

Schließlich heißt K ein archimedisch angeordneter Körper, wenn zusätzlich das Archimedische Axiom gilt:

A3 Zu jedem a gibt es ein $n \in \mathbb{N}$ mit $a < n$.

Dabei bezeichnet n das Element, das induktiv definiert ist durch $n = (n - 1) + 1$.

Damit ist also $\mathbb{Q}$ ein angeordneter Körper.

Bemerkung. Das archimedische Axiom ist von den übrigen Axiomen unabhängig, denn es gibt Körper, die alle Axiome bis auf **A3** erfüllen.

3. Reelle Zahlen

Die natürlichen Zahlen und die ganzen Zahlen sind in dem Sinne unvollständig, dass in ihnen die Subtraktion und Division nicht uneingeschränkt möglich sind. In $\mathbb{Q}$ ist das anders, aber dennoch ist $\mathbb{Q}$ als Zahlenbereich nicht ausreichend, um in der Natur vorkommende Größen zu beschreiben: man denke nur an die Länge der Diagonalen im Einheitsquadrat, die keine rationale Zahl sein kann (siehe das Beispiel 2 auf Seite 6).

Dieses Problem ließe sich noch algebraisch lösen, durch eine ähnliche Konstruktion wie oben: Betrachte die Menge $\mathbb{Q} \times \mathbb{Q}$ und definiere darauf Addition und Multiplikation durch

$$\begin{aligned}(m_1, m_2) + (n_1, n_2) &= (m_1 + n_1, m_2 + n_2) \\ (m_1, m_2) \cdot (n_1, n_2) &= (m_1 n_1 + 2m_2 n_2, m_1 n_2 + n_1 m_2)\end{aligned}$$

Man überzeugt sich schnell davon, dass dadurch wieder ein Körper gegeben ist (Axiome **K1** bis **K5**), in dem die Gleichung $x^2 = 2$ eine Lösung hat, nämlich $x = (0, 1)$.

Wie kommt man auf diese Multiplikationsregel? Indem man ein Paar (a, b) als $a + b\sqrt{2}$ auffasst, denn es gilt

$$(m_1 + m_2\sqrt{2})(n_1 + n_2\sqrt{2}) = (m_1 n_1 + 2m_2 n_2) + (m_1 n_2 + m_2 n_1)\sqrt{2}.$$

Aber selbst in dieser Erweiterung finden wir wieder Gleichungen, die wir nicht lösen können, etwa $x^2 = 3$. Nun ist es zwar möglich, auf rein algebraische Weise einen Körper zu konstruieren, in dem jede algebraische Gleichung eine Lösung hat (dies nennt man dann den algebraischen Abschluss), aber selbst darin gibt es keine Zahl, die dem Umfang des Einheitskreises entspricht — die Quadratur des Kreises ist nicht möglich. Andererseits können wir diese Zahl (2π) beliebig genau durch Brüche approximieren.

Abhilfe liefern die **reellen Zahlen**. Es gibt verschiedene Möglichkeiten, diese zu konstruieren; mit am gebräuchlichsten ist die Methode der **Dedekindschen Schnitte**¹. Diese Konstruktion soll hier nur skizziert werden; sie gehört eigentlich in eine Analysis-Vorlesung.

Definition 3.1. Ein (Dedekind-)Schnitt α ist eine Teilmenge $\alpha \subset \mathbb{Q}$ mit folgenden Eigenschaften:

- (1) $\alpha \neq \emptyset$ und $\alpha \neq \mathbb{Q}$.
- (2) Zu jedem $x \in \alpha$ gibt es ein $y \in \alpha$ mit $x < y$.
- (3) Ist $x \in \alpha$ und $y < x$, so ist auch $y \in \alpha$.

Eigenschaft (2) heißt, dass die Menge α kein größtes Element hat.

Beispiele. 1. Jede rationale Zahl r bestimmt einen Schnitt α_r durch

$$\alpha_r = \{x \in \mathbb{Q} \mid x < r\}.$$

Einen solchen Schnitt nennt man von der rationalen Zahl r **erzeugt**. Zwei rationale Zahlen r, s erzeugen offenbar genau dann den selben Schnitt, wenn $r = s$ ist. Vermöge der Zuordnung $r \mapsto \alpha_r$ kann man also $\mathbb{Q}$ als Teilmenge der Menge aller Schnitte auffassen.

2. Die Menge

$$M = \{x \in \mathbb{Q} \mid x < 0 \vee x^2 < 2\}$$

ist ebenfalls ein Schnitt: Zu jeder rationalen Zahl x mit $x^2 < 2$ gibt es eine größere y , die ebenfalls noch kleiner als 2 ist, nämlich $\frac{x+2}{2}$. Die Eigenschaften (1) und (3) sind klar. M ist ein Schnitt, der nicht von der Form des ersten Beispiels ist, d.h. es gibt keine rationale Zahl r mit $M = \alpha_r$.

Bemerkung. In der Literatur wird als Dedekindscher Schnitt oft ein Paar (A, B) von Teilmengen von $\mathbb{Q}$ bezeichnet, so dass (i) $A \cup B = \mathbb{Q}$, (ii) jedes Element von A kleiner als jedes Element von B ist, und (iii) A kein größtes Element besitzt. Dies ist zu der Definition oben äquivalent: ist α ein Schnitt im Sinne der Definition, so ist $(\alpha, \mathbb{Q} - \alpha)$ ein Paar, das (i) – (iii) erfüllt, und ist (A, B) ein Paar mit diesen Eigenschaften, so ist A ein Schnitt wie in der Definition. (Es gibt noch Varianten: gelegentlich wird auch die Bedingung (iii) ersetzt durch „ A hat kein grösstes oder B kein kleinstes Element“ oder sogar ganz weggelassen, was dann aber dazu führt, dass rationalen Zahlen zwei mögliche Paare entsprechen. Die hier verwendete Definition erspart lästige Fallunterscheidungen).

Man definiert nun $\mathbb{R}$ als die Menge der Schnitte:

Definition 3.2. Eine reelle Zahl ist in Dedekindscher Schnitt. Sei $\mathbb{R}$ die Menge der reellen Zahlen.

Auf $\mathbb{R}$ definiert man eine Ordnungsrelation durch Inklusion von Teilmengen: Für $\alpha, \beta \in \mathbb{R}$ sei

$$\alpha \leq \beta \quad \text{genau dann, wenn} \quad \alpha \subset \beta.$$

Man schreibt dann wieder $\alpha < \beta$, falls $\alpha \leq \beta$ und $\alpha \neq \beta$ ist. Diese Ordnung ist wegen Eigenschaft (3) der Schnitte eine Totalordnung.

Eine reelle Zahl heißt **irrational**, wenn $\mathbb{Q} - \alpha$ kein kleinstes Element enthält. Dies ist genau dann der Fall, wenn α nicht von der Form α_r für ein $r \in \mathbb{Q}$, also nicht rational ist.

Die wesentliche Eigenschaft der reellen Zahlen (im Gegensatz zu den rationalen Zahlen) beschreibt der nun folgende Satz. Dabei nennen wir eine Teilmenge $A \subset \mathbb{R}$ nach oben (nach unten) beschränkt, wenn es ein $\beta \in \mathbb{R}$ gibt, so dass $\alpha \leq \beta$ (bzw. $\beta \leq \alpha$) für jedes $\alpha \in A$ gilt. In diesem Fall nennt man β eine obere (eine untere) Schranke für A .

Satz 3.3. Jede nichtleere, nach oben beschränkte Teilmenge $A \subset \mathbb{R}$ hat eine kleinste obere Schranke.

¹Richard Dedekind, deutscher Mathematiker, 1831–1916. Seine Konstruktion der reellen Zahlen veröffentlichte er 1872 in der Schrift „Stetigkeit und irrationale Zahlen“.

Beweis. Sei $A \subset \mathbb{R}$ nichtleer, und sei $\gamma \in \mathbb{R}$ eine obere Schranke für A . Definiere die Menge

$$\sup A := \bigcup_{\alpha \in A} \alpha.$$

Es ist zunächst zu zeigen, dass $\sup A$ eine reelle Zahl ist, also die Eigenschaften (1) – (3) eines Dedekind-Schnitts erfüllt.

- (1) $\sup A \neq \emptyset$ ist klar, denn A enthält wenigstens eine nichtleere Menge $\alpha \subset \mathbb{Q}$, und wegen $\alpha \subset \sup A$ per Konstruktion ist dann auch $\sup A$ nicht leer. Da γ ein Schnitt ist, enthält $\mathbb{Q} - \gamma$ wenigstens eine rationale Zahl x , und da γ eine obere Schranke ist, ist jedes $\alpha \in A$ eine Teilmenge von γ . Folglich ist x in keiner der Mengen α enthalten und damit auch nicht in $\sup A$. Dies zeigt $\sup A \neq \mathbb{Q}$.
- (2) Sei $x \in \sup A$. Dann ist x in einer der Mengen α enthalten. Weil α die Eigenschaft (2) hat, gibt es ein $y > x$ mit $y \in A \subset \sup A$.
- (3) Sei $x \in \sup A$ und $y < x$. Dann ist $x \in \alpha$ für ein $\alpha \in A$, und da α die Eigenschaft (3) hat, ist auch jedes kleinere y in α , mithin in $\sup A$.

Also ist $\sup A$ eine reelle Zahl; per Konstruktion ist $\sup A$ eine oberer Schranke für A . Zu zeigen bleibt, dass $\sup A$ eine kleinste obere Schranke ist, d.h. dass für jede obere Schranke β von A gilt: $\alpha \leq \beta$.

Sei also β eine beliebige obere Schranke für A und $x \in \sup A$. Dann ist $x \in \alpha$ für ein $\alpha \in A$, und wegen $\alpha \subset \beta$ (obere Schranke) ist $x \in \beta$. Dies zeigt $\sup A \subset \beta$, also $\sup A \leq \beta$. $\square$

Kleinste obere Schranken sind eindeutig bestimmt: sind β und β' beides kleinste obere Schranken, so muss $\beta \leq \beta'$ und $\beta' \leq \beta$ gelten, also $\beta = \beta'$. Man kann also von *der* kleinsten oberen Schranke sprechen. Die kleinste obere Schranke $\sup A$ einer nach oben beschränkten Menge A nennt man das **Supremum** von A .

Bemerkung. Es gilt die analoge Aussage für nach unten beschränkte Teilmengen von $\mathbb{R}$: Jede nach unten beschränkte Teilmenge $B \subset \mathbb{R}$ hat eine größte untere Schranke. Diese nennt man das **Infimum** von B , geschrieben $\inf B$.

Die in dem Satz formulierte Eigenschaft der reellen Zahlen nennt man **Vollständigkeit**².

In Analogie zu Definition 3.1 nennt man eine Teilmenge $A \subset \mathbb{R}$ einen Schnitt von $\mathbb{R}$, wenn sie die Eigenschaften (1) – (3) dieser Definition erfüllt. Eine solche Teilmenge ist insbesondere nach oben beschränkt: hätte A keine obere Schranke, so gäbe es zu jedem $\beta \in \mathbb{R}$ ein $\alpha \in A$ mit $\beta \leq \alpha$, so dass wegen (3) $\beta \in A$ folgte; es wäre also $A = \mathbb{R}$ im Widerspruch zu (1). Nach dem eben bewiesenen Satz hat also jeder Schnitt $A \subset \mathbb{R}$ eine kleinste obere Schranke $\sup A \in \mathbb{R}$, die aber wegen (2) nicht in A liegt. Sei nun

$$X = \{\xi \in \mathbb{R} \mid \xi < \sup A\},$$

X ist also von der reellen Zahl $\sup A$ erzeugt. (Wie früher sind zwei Schnitte Y, Z , die durch reelle Zahlen η, ζ erzeugt sind, genau dann gleich, wenn η und ζ gleich sind; so kann man wieder $\mathbb{R}$ als Teilmenge der Menge aller Schnitte auffassen.) Dann gilt $X \subset A$, denn es gilt $\alpha \leq \sup A$ und $\alpha \neq \sup A$, also $\alpha < \sup A$ für jedes $\alpha \in A$. Aber es ist auch $A \subset X$, denn jedes Element η , das nicht in A liegt, ist eine obere Schranke für A , muss also $\geq \sup A$ sein und kann daher auch nicht in A liegen. Es ist also $A = X$, d.h. A ist durch die reelle Zahl $\sup A$ erzeugt. Des bedeutet, dass die Wiederholung der Dedekindschen Konstruktion nicht zu einer Erweiterung der Menge $\mathbb{R}$ führt, oder anders ausgedrückt:

Korollar 3.4 (Vollständigkeitssatz). $\mathbb{R}$ ist Dedekind-vollständig: Jeder Schnitt in $\mathbb{R}$ ist von einer reellen Zahl erzeugt. $\square$

Die reellen Zahlen sind nun als total geordnete Menge beschrieben; es fehlen noch die arithmetischen Operationen.

²In der Analysis wird man eine andere Charakterisierung der Vollständigkeit (durch sogenannte Cauchy-Folgen) kennenlernen.

Definition 3.5. (a) Die Elemente $\mathbf{0} \in \mathbb{R}$ und $\mathbf{1} \in \mathbb{R}$ seien gegeben durch

$$\mathbf{0} := \{x \in \mathbb{Q} \mid x < 0\} = \alpha_0$$

$$\mathbf{1} := \{x \in \mathbb{Q} \mid x < 1\} = \alpha_1$$

(b) Für $\alpha, \beta \in \mathbb{R}$ sei

$$\alpha + \beta := \{x + y \mid x \in \alpha, y \in \beta\}$$

$$-\alpha := \{x \in \mathbb{Q} \mid -x \notin \alpha \wedge (\exists y \in \alpha \ y < -x)\}$$

$$|\alpha| := \begin{cases} \alpha & \text{falls } 0 \leq \alpha, \\ -\alpha & \text{falls } \alpha < 0. \end{cases}$$

(c) Für $\alpha, \beta > 0$ sei

$$\alpha \cdot \beta := \{z \in \mathbb{Q} \mid z \leq 0 \vee (\exists x \in \alpha \ \exists y \in \beta \ (z = xy \wedge xy > 0))\}$$

sowie allgemein

$$\alpha \cdot \beta := \begin{cases} 0 & \text{falls } \alpha = 0 \text{ oder } \beta = 0, \\ |\alpha| \cdot |\beta| & \text{falls } \alpha > 0, \beta > 0 \text{ oder } \alpha < 0, \beta < 0, \\ -(|\alpha| \cdot |\beta|) & \text{sonst.} \end{cases}$$

(d) Für $\alpha > 0$ sei

$$\alpha^{-1} := \left\{ x \in \mathbb{Q} \mid \begin{array}{l} x \leq 0 \text{ oder } (x > 0 \text{ und } \frac{1}{x} \notin \alpha \text{ und } \frac{1}{x} \text{ ist} \\ \text{nicht das kleinste Element von } \mathbb{Q} - \alpha). \end{array} \right\}$$

und für $\alpha < 0$ sei $\alpha^{-1} = -(|\alpha|)^{-1}$.

Nun muss man die Körperaxiome nachrechnen.

Die Assoziativ- und Kommutativgesetze folgen leicht aus den entsprechenden Gesetzen für rationale Zahlen, ebenso das Distributivgesetz. Für die Axiome **K4** und **K5** muss man etwas mehr arbeiten; dies sei hier exemplarisch für die Addition durchgeführt.

Sei $\alpha \in \mathbb{R}$.

K4: Per Definition ist

$$\alpha + \mathbf{0} = \{x + y \mid x \in \alpha \wedge y < 0\},$$

zu zeigen sind die beiden Inklusionen $\alpha + \mathbf{0} \subset \alpha$ und $\alpha \subset \alpha + \mathbf{0}$.

Die Inklusion $\alpha + \mathbf{0} \subset \alpha$ ist klar, denn jedes Element $x + y$ von $\alpha + \mathbf{0}$ mit $x \in \alpha, y \in \mathbf{0}$ ist kleiner als x , also nach Eigenschaft (3) wieder in α . Sei umgekehrt $z \in \alpha$. Dann gibt es (Eigenschaft (2)) ein $x \in \alpha$ mit $z < x$. Setze $y = z - x$, dann ist $y \in \mathbf{0}$, also $z = x + y \in \alpha + \mathbf{0}$.

K5: Zu zeigen sind die Inklusionen $\alpha + (-\alpha) \subset \mathbf{0}$ und $\mathbf{0} \subset \alpha + (-\alpha)$. Die erste ist wieder leicht: Sei $x \in \alpha + (-\alpha)$, also $x = u + v$ mit $u \in \alpha$ und $-v \in (-\alpha)$, d.h. $-v \in \mathbb{Q} - \alpha$ und nicht kleinstes Element dieser Menge. Es folgt $u < -v$, also $u + v < 0$, mithin $u + v \in \mathbf{0}$.

Für die umgekehrte Inklusion sei $x \in \mathbf{0}$, d.h. x eine negative rationale Zahl. Wegen $\emptyset \neq \alpha \neq \mathbb{Q}$ gibt es ein $y \in \mathbb{Q} - \alpha$ und ein $z \in \alpha$. Betrachte die Menge

$$M = \{m \in \mathbb{N} \mid y + mx \in \alpha\}.$$

M ist nicht leer, denn für jedes $m > \frac{z-y}{x}$ gilt $y + mx < y + \frac{z-x}{x}x = z$ (es ist $x < 0$, daher dreht sich die Ungleichung um), also $y + mx \in \alpha$ nach Eigenschaft (3). Folglich hat M ein kleinstes Element n . Dann gilt

$$y + nx \in \alpha \quad \text{und} \quad y + (n-1)x \in \mathbb{Q} - \alpha.$$

Man kann aber nicht ausschließen, dass $y + (n-1)x$ das kleinste Element von $\mathbb{Q} - \alpha$ ist. Wenn dies der Fall sein sollte, ersetze y durch $y' = y + \frac{x}{2}$; dann gilt für das gleiche n

$$y' + nx \in \alpha \quad \text{und} \quad y' + (n-1)x \in \mathbb{Q} - \alpha.$$

und $y' + (n-1)x$ ist kein kleinstes Element von $\mathbb{Q} - \alpha$ mehr, denn es ist größer als $y + (n-1)x$. Setzt man nun $x_0 = y + (n-1)x$ bzw. $x_0 = y' + (n-1)x$, so gilt $-x_0 \in (-\alpha)$ und $x_0 + x \in \alpha$, also

$$x = (x_0 + x) + (-x_0) \in \alpha + (-\alpha).$$

Die Argumente für die Multiplikation sind ähnlich.

Um zu der aus der Schule gewohnten Dezimalbruchentwicklung einer reellen Zahl zu gelangen, gibt es ein einfaches Verfahren. Sei $\alpha \in \mathbb{R}$, zunächst mit $\alpha > 0$. Sei die Folge $(x_k)_{k \in \mathbb{N}}$ induktiv definiert durch

$$x_0 = \max\{n \in \mathbb{N} \mid n \in \alpha\}$$

wobei max für das Maximum stehe, dies ist der ganze Anteil von α , und dann

$$x_1 = \max\{m \in \{0, 1, \dots, 9\} \mid x_0 + \frac{m}{10} \in \alpha\}$$

$$x_2 = \max\{m \in \{0, 1, \dots, 9\} \mid x_0 + \frac{x_1}{10} + \frac{m}{100} \in \alpha\}$$

$$\vdots$$

$$x_k = \max\{m \in \{0, 1, \dots, 9\} \mid x_0 + \frac{x_1}{10} + \frac{x_2}{100} + \dots + \frac{x_{k-1}}{10^{k-1}} + \frac{m}{10^k} \in \alpha\}$$

Dann ist

$$x_0, x_1 x_2 x_3 \dots x_k \dots$$

die Dezimalbruchentwicklung von α . Für negative Zahlen leitet man daraus leicht ein entsprechendes Verfahren ab (Übungsaufgabe).

Ganz analog kann man reelle Zahlen auch in andere als Dezimalbrüche entwickeln, etwa als sogenannte **Dualzahl** mit den Ziffern 0 und 1. Sei β eine reelle Zahl, ohne Einschränkung sei $0 < \beta < 1$. Setze dann

$$y_1 = \begin{cases} 1 & \text{falls } \frac{1}{2} \in \beta, \\ 0 & \text{sonst,} \end{cases}$$

$$y_2 = \begin{cases} 1 & \text{falls } \frac{y_1}{2} + \frac{1}{4} \in \beta, \\ 0 & \text{sonst,} \end{cases}$$

$$\vdots$$

$$y_k = \begin{cases} 1 & \text{falls } (\sum_{i=1}^{k-1} y_i 2^{-i}) + 2^{-k} \in \beta, \\ 0 & \text{sonst.} \end{cases}$$

4. Komplexe Zahlen

Es gibt keine reellen Zahlen, die Gleichungen wie

$$z^2 + 1 = 0$$

$$z^2 + z + 1 = 0$$

lösen, wie man sofort sieht, wenn man die bekannte Formel für quadratische Gleichungen benutzen möchte, denn es treten Wurzeln aus negativen Zahlen auf. Um Abhilfe zu schaffen, konstruiert man eine Erweiterung der reellen Zahlen, die solche Lösungen enthält.

Betrachte dazu $\mathbb{R} \times \mathbb{R}$ mit komponentenweiser Addition, also

$$(x_1, y_1) + (x_2, y_2) := (x_1 + x_2, y_1 + y_2),$$

und folgender Multiplikation:

$$(x_1, y_1) \cdot (x_2, y_2) := (x_1 x_2 - y_1 y_2, x_1 y_2 + x_2 y_1).$$

Diese Multiplikation ist assoziativ:

$$\begin{aligned} ((x_1, y_1) \cdot (x_2, y_2)) \cdot (x_3, y_3) &= (x_1x_2 - y_1y_2, x_1y_2 + x_2y_1) \cdot (x_3, y_3) \\ &= (x_1x_2x_3 - x_3y_1y_2 - x_1y_2y_3 - x_2y_1y_3, x_1x_2y_3 + x_1x_3y_2 + x_2x_3y_1 - y_1y_2y_3) \\ &= (x_1, y_1) \cdot (x_2x_3 - y_2y_3, x_2y_3 + x_3y_2) \\ &= (x_1, y_1) \cdot ((x_2, y_2) \cdot (x_3, y_3)) \end{aligned}$$

Noch leichter sieht man, dass sie auch kommutativ ist, denn der definierende Ausdruck bleibt gleich, wenn man x_1 mit x_2 und y_1 mit y_2 vertauscht. Genauso einfach lässt sich das Distributivgesetz nachrechnen. Für die Addition ist offenbar $(0, 0)$ neutrales Element, und für die Multiplikation ist es $(1, 0)$, denn

$$(1, 0) \cdot (x, y) = (1 \cdot x - 0 \cdot y, 1 \cdot y + 0 \cdot x) = (x, y).$$

Additives Inverses zu (x, y) ist dann $(-x, -y)$, und das multiplikative Inverse zu $(x, y) \neq (0, 0)$ ist

$$(x, y)^{-1} = \left(\frac{x}{x^2 + y^2}, \frac{-y}{x^2 + y^2} \right),$$

wie man aus der Rechnung

$$(x, y) \cdot \left(\frac{x}{x^2 + y^2}, \frac{-y}{x^2 + y^2} \right) = \left(\frac{x^2 - y(-y)}{x^2 + y^2}, \frac{xy + x(-y)}{x^2 + y^2} \right) = (1, 0)$$

sieht.

Damit ist nachgewiesen, dass die Menge $\mathbb{R} \times \mathbb{R}$ mit dieser Addition und Multiplikation ein Körper ist, den man den Körper der **komplexen Zahlen** nennt und mit $\mathbb{C}$ bezeichnet. $\mathbb{R}$ ist dabei die Teilmenge aller Zahlen mit zweiter Komponente Null; mit anderen Worten, die Abbildung

$$\iota: \mathbb{R} \longrightarrow \mathbb{C}, \quad x \mapsto (x, 0)$$

ist injektiv, und sie erhält auch die arithmetischen Operationen. Man sagt daher, dass $\mathbb{R}$ ein **Unter-** oder **Teilkörper** von $\mathbb{C}$ sei.

In $\mathbb{C}$ hat die Gleichung $z^2 + 1 = 0$ eine Lösung: es gilt

$$(0, 1) \cdot (0, 1) = (-1, 0),$$

und letztere Zahl ist die reelle Zahl -1 (wobei man reelle Zahlen mit ihrem Bild unter ι identifiziert). Diese Zahl $(-1, 0)$ wird üblicherweise mit i bezeichnet; es ist also

$$i := (0, 1) \quad \text{mit} \quad i^2 = -1;$$

man nennt i auch die **imaginäre Einheit**. Jede komplexe Zahl (x, y) kann man schreiben als

$$(x, y) = (x, 0) \cdot (1, 0) + (y, 0) \cdot (0, 1) = x \cdot 1 + y \cdot i = x + iy.$$

In Zukunft sollen, wie gemeinhin üblich, komplexe Zahlen in dieser Form geschrieben werden, und nicht mehr als Paare. Diese Schreibweise ist insofern gerechtfertigt, weil sie ihrerseits die Rechengesetze erweitert: für $z_1 = (x_1, y_1)$ und $z_2 = (x_2, y_2)$ gilt nämlich

$$\begin{aligned} (x_1 + iy_1) \cdot (x_2 + iy_2) &= x_1x_2 + x_1(iy_2) + (iy_1)x_2 + (iy_1)(iy_2) \\ &= x_1x_2 + i^2y_1y_2 + i(x_1y_2 + x_2y_1) \\ &= x_1x_2 - y_1y_2 + i(x_1y_2 + x_2y_1) = (x_1, y_1) \cdot (x_2, y_2) = z_1 \cdot z_2. \end{aligned}$$

Was bisher gezeigt wurde, sei noch einmal zu einem Satz zusammengefasst:

Satz 4.1. *Es gibt einen Körper $(\mathbb{C}, +, \cdot, 0, 1)$, so dass gilt:*

- (1) $\mathbb{C}$ enthält $\mathbb{R}$ als Teilkörper.
- (2) Es gibt ein $i \in \mathbb{C}$ mit $i^2 = -1$.
- (3) Jedes $z \in \mathbb{C}$ lässt sich eindeutig darstellen als $z = x + iy$ mit $x, y \in \mathbb{R}$.

□

Für eine komplexe Zahl $z = x + iy$ nennt man

$\operatorname{Re}(z) := x$ den **Realteil** von z ,

$\operatorname{Im}(z) := y$ den **Imaginärteil** von z .

Auf den komplexen Zahlen lässt sich keine Ordnung definieren, die den Anordnungsaxiomen **A1** und **A2** genügt: Setzt man etwa $0 < i$, so folgt auch $0 < i^2 = -1$, und damit auch $0 < (-1) \cdot i = -i$, also $0 < i + (-i) = 0$, was absurd ist. Auf gleiche Weise führt die Annahme $i < 0$ zu einem Widerspruch, so dass man also 0 und i nicht vergleichen kann. Mit anderen Worten: es gibt keine Ordnungsrelation, die $\mathbb{C}$ zu einem angeordneten Körper macht.

Die Definition als reelle Zahlenpaare ermöglicht es, komplexe Zahlen geometrisch als Punkte in der Ebene $\mathbb{R}^2$ bzw. als Vektoren in der Ebene zu veranschaulichen: die komplexe Zahl $z = x + iy$ entspricht dabei dem Punkt mit reellen Koordinaten (x, y) .

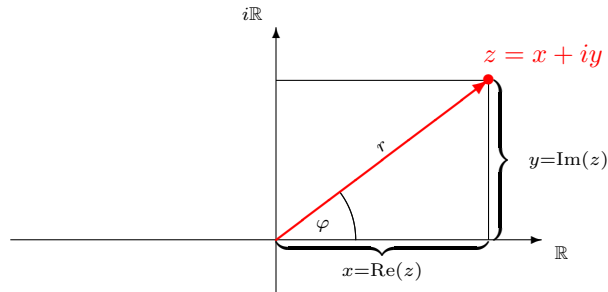


ABBILDUNG 1. Geometrische Veranschaulichung komplexer Zahlen

Die Länge dieses Vektors, also der Abstand des Punktes (x, y) vom Ursprung, beträgt laut Pythagoras $r = \sqrt{x^2 + y^2}$. Man definiert daher den **Betrag** der komplexen Zahl $z = x + iy$ als

$$|z| := \sqrt{x^2 + y^2}.$$

Dann ist $|z|$ eine nichtnegative reelle Zahl, und es gilt $|z| = 0$ genau dann, wenn $z = 0$ ist. Weiterhin gilt für $z_1 = x_1 + iy_1$ und $z_2 = x_2 + iy_2$:

$$\begin{aligned} |z_1|^2 \cdot |z_2|^2 &= (x_1^2 + y_1^2) \cdot (x_2^2 + y_2^2) \\ &= x_1^2 x_2^2 + y_1^2 y_2^2 + x_1^2 y_2^2 + x_2^2 y_1^2 \\ &= x_1^2 x_2^2 - 2x_1 x_2 y_1 y_2 + y_1^2 y_2^2 + x_1^2 y_2^2 + 2x_1 y_2 x_2 y_1 + x_2^2 y_1^2 \\ &= (x_1 x_2 - y_1 y_2)^2 + (x_1 y_2 + x_2 y_1)^2 \\ &= |z_1 \cdot z_2|^2, \end{aligned}$$

also

$$|z_1 \cdot z_2| = |z_1| \cdot |z_2|.$$

Da man einen Punkt der Ebene genauso gut durch die Angabe seines Abstandes vom Ursprung zusammen mit dem Winkel der Geraden durch den Punkt und den Ursprung zur x -Achse beschreiben kann, gilt das gleiche für komplexe Zahlen. Ist also $r \in \mathbb{R}$ eine nichtnegative Zahl und φ ein Winkel, $0 \leq \varphi < 2\pi$, so definiert das Paar (r, φ) die komplexe Zahl

$$z = r \cos \varphi + ir \sin \varphi.$$

Man nennt die Koordinaten r und φ auch **Polarkoordinaten**. In dieser Darstellung ist die Multiplikation einfacher: seien $z_1 = (r_1, \varphi_1)$ und $z_2 = (r_2, \varphi_2)$ zwei komplexe Zahlen in Polarkoordinaten, dann gilt nach den Additionstheoremen für die Winkelfunktionen

$$\begin{aligned}
z_1 \cdot z_2 &= (r_1 \cos \varphi_1 + ir_1 \sin \varphi_1) \cdot (r_2 \cos \varphi_2 + ir_2 \sin \varphi_2) \\
&= r_1 r_2 (\cos \varphi_1 \cos \varphi_2 - \sin \varphi_1 \sin \varphi_2) + ir_1 r_2 (\cos \varphi_1 \sin \varphi_2 + \sin \varphi_1 \cos \varphi_2) \\
&= r_1 r_2 \cos(\varphi_1 + \varphi_2) + i \sin(\varphi_1 + \varphi_2)
\end{aligned}$$

also wieder zurückübersetzt in Polarkoordinaten

$$(r_1 \varphi_1) \cdot (r_2 \varphi_2) = (r_1 r_2, \varphi_1 + \varphi_2).$$

Mit anderen Worten: die Längen werden multipliziert und die Winkel addiert.

Den Betrag einer komplexen Zahl kann man noch anders schreiben. Die Abbildung $\mathbb{C} \rightarrow \mathbb{C}$, die durch die Vorschrift

$$z = x + iy \mapsto \bar{z} := x - iy$$

definiert ist, bezeichnet man als **komplexe Konjugation**. Geometrisch bedeutet dies die Spiegelung an der reellen Achse. Dann gilt

$$z \cdot \bar{z} = |z|^2.$$

Offenbar gilt auch $|z| = |\bar{z}|$ für jedes $z \in \mathbb{C}$.

Weitere Rechenregeln verifiziert man leicht aus den Definitionen:

Lemma 4.2. Für $z, w \in \mathbb{C}$ gilt

- (a) $\overline{z + w} = \bar{z} + \bar{w}$, $\overline{z \cdot w} = \bar{z} \cdot \bar{w}$.
- (b) $z + \bar{z} = 2 \operatorname{Re}(z)$, $z - \bar{z} = 2i \cdot \operatorname{Im}(z)$.
- (c) $|\operatorname{Re}(z)| \leq |z|$, $|\operatorname{Im}(z)| \leq |z|$.

□

Geometrisch klar (wenn auch eines formalen Beweises bedürftig) ist das folgende Resultat:

Satz 4.3 (Dreiecksungleichung). Seien $z, w \in \mathbb{C}$. Dann gilt:

$$|z + w| \leq |z| + |w|.$$

Beweis. Wegen Lemma 4.2 (c) gilt

$$z\bar{w} + \bar{z}w = z\bar{w} + \overline{z\bar{w}} \leq 2|z\bar{w}| = 2|z||w|$$

und damit

$$|z + w|^2 = |z|^2 + z\bar{w} + \bar{z}w + |w|^2 \leq |z|^2 + 2|z||w| + |w|^2 = (|z| + |w|)^2$$

woraus die Behauptung folgt. □

Zu Beginn dieses Abschnitts wurde die Einführung der komplexen Zahlen damit motiviert, dass man in $\mathbb{C}$ jede quadratische Gleichung lösen kann. Es gilt aber noch viel mehr:

Satz 4.4 (Fundamentalsatz der Algebra). Jedes Polynom mit komplexen Koeffizienten hat eine komplexe Nullstelle.

Ein Beweis dieses Satzes geht weit über das hinaus, was in einem Vorkurs behandelt werden kann. Interessant ist aber eine Folgerung, die man aus dem Satz ziehen kann. Ist

$$p(z) = a_n z^n + a_{n-1} z^{n-1} + \cdots + a_1 z_1 + a_0, \quad a_i \in \mathbb{C}, \quad a_n \neq 0$$

ein komplexes Polynom vom Grad n und z_1 eine Nullstelle dieses Polynoms, so ist $p(z)$ durch $(z - z_1)$ teilbar, d.h. es gibt ein Polynom $q(z)$ mit $p(z) = q(z) \cdot (z - z_1)$. Aber dann ist $q(z)$ ebenfalls ein komplexes Polynom, hat also eine komplexe Nullstelle, usw. Induktiv folgt schließlich, dass es komplexe Zahlen $z_1, z_2, \dots, z_n$ gibt (die nicht verschieden sein müssen), so dass

$$p(z) = (z - z_1) \cdot (z - z_2) \cdot \cdots \cdot (z - z_n)$$

gilt. Man sagt auch, $p(z)$ zerfalle in Linearfaktoren.