

Lehrbuch der Mathematik

ZWÖLFTES SCHULJAHR

LEHRBUCH
DER
MATHEMATIK

FÜR DIE OBERSCHULE

12. SCHULJAHR

AUSGABE 1954



VOLK UND WISSEN VOLKSEIGENER VERLAG BERLIN

1955

Der Teil A wurde von Dr. Gustav Beyrodt,
der Teil B von Dr. Johannes Kusch, der Teil C von Edgar Brey
und der Teil D von Dr. Hanns Joachim Weinert verfaßt.

Zeichnungen: Kurt Dornbusch

Redaktionelle Bearbeitung: Johannes Riedel, Dietrich Kind und Rosemarie Acklow

Die erste Auflage erschien in zwei Teilen

Redaktionseschluß: 30. April 1954

Bestell-Nr. 00913-2 · 3,65 DM · Lizenz Nr. 203 · 1000-P-005515 (DN)

Satz und Druck: Druckhaus „Maxim Gorki“, Altenburg

A. Fehlerrechnung und Näherungslösungen

I. Fehlerrechnung

In den Naturwissenschaften und in der Technik haben wir es mit Größen zu tun, die z. B. als Längen, Massen, Kräfte oder Zeiten gemessen werden oder die aus solchen Messungen berechnet sind. Die berechneten Größen sind dann Funktionen der unmittelbar gemessenen Größen.

Die Erfahrung hat gezeigt, daß bei wiederholten Messungen des gleichen Gegenstandes die Ergebnisse der Einzelmessungen voneinander abweichen, auch wenn alle Messungen mit der gleichen Sorgfalt, mit dem gleichen Instrument und unter möglichst gleichen Voraussetzungen durchgeführt wurden. Wollen wir also eine Größe, die gemessen werden muß, bestimmen, so müssen wir die Gesamtheit aller Einzelmessungen erfassen und mathematisch auswerten.

Im Lehrbuch der Physik für das 9. Schuljahr ist schon eine Einführung in die Meßkunde gegeben worden. Auf den Seiten 11 bis 13 ist dabei die Entstehung von Meßfehlern und ihre Berücksichtigung besprochen worden. Die Untersuchungen beschränken sich dort auf die Grundlagen, die bei der Auswertung physikalischer Experimente unbedingt gebraucht werden. Auch im Lehrbuch der Mathematik für das 11. Schuljahr ist auf den Seiten 47 bis 49 ein einzelnes Problem behandelt worden. Wir müssen nun untersuchen, welchen Einfluß die Fehler der Ausgangswerte auf die Ergebnisse von Berechnungen haben.

1. Grundlegende Begriffe

Die zu messende Größe, z. B. eine Kraft, eine Länge, eine Zeit, bezeichnet man als **Meßgröße**.

Der Gegenstand, an dem die Messung ausgeführt wird, heißt der **Meßgegenstand** oder das **Meßobjekt**.

Den Wert, den die Messung ergibt, bezeichnen wir als **Meßwert**. Die Differenz zwischen dem Meßwert und dem richtigen Wert ist der **Fehler**.

Welchen Wert wir als den richtigen Wert bezeichnen, werden wir im Abschnitt 2 feststellen. Der Fehler ist positiv, wenn der Meßwert größer als der richtige Wert ist, und negativ, wenn der Meßwert kleiner als der richtige Wert ist.

Fehler sind entweder systematischer oder zufälliger Natur.

Zu den **systematischen Fehlern** gehören Fehler der benutzten Meßgeräte und die Fehler, die ihren Grund in meßbaren Umwelteinflüssen haben. Meßbare

Umwelteinflüsse sind z. B. Änderungen der Temperatur, der Luftfeuchtigkeit und des Luftdruckes. Sie haben eine bestimmte Größe und ein bestimmtes Vorzeichen. Systematische Fehler lassen sich im allgemeinen mathematisch erfassen und bei jeder Einzelmessung korrigieren.

Auch nach Ausschaltung der systematischen Fehler treten trotz sorgfältigster Messung und Verwendung von Präzisionsinstrumenten noch Fehler auf, die bei wiederholten Messungen nach Größe und Vorzeichen schwanken. Diese Fehler bezeichnen wir als **zufällige Fehler**. Sie lassen sich durch Berichtigungen nicht ausschalten. Wir werden uns in den folgenden Abschnitten nur mit den zufälligen Fehlern beschäftigen.

2. Mittelwert und durchschnittlicher Fehler

Da wir wissen, daß jeder Meßwert mit Ungenauigkeiten behaftet ist, verlassen wir uns grundsätzlich nicht auf eine Messung, sondern wir bestimmen Meßgrößen durch einen Satz von Messungen unter möglichst gleichen Voraussetzungen. Die Zahl der Messungen ist von der geforderten Genauigkeit abhängig. So werden bei der Landesvermessung und beim Feldmessen Winkelmessungen mindestens in drei Sätzen zu je zwei Messungen, unter Umständen sogar in zwölf Sätzen, und Längenmessungen je dreimal in zwei Richtungen durchgeführt. Aus diesen Meßwerten wird das arithmetische Mittel bestimmt. Den auf diese Weise ermittelten Wert kann man als einen besseren Wert für die wirkliche Größe des Meßobjektes ansehen als den durch eine Messung bestimmten Wert. Da wir die wirkliche Größe des Meßobjektes nicht kennen, wollen wir ihn als den „richtigen“ Wert annehmen.

Beispiel 1:

Bei einem Streckenzug von P_1 nach P_2 werden in je 3 Messungen mit Theodolit und Meßlatte von P_1 aus die Entfernungen 235,73 m, 235,69 m, 235,76 m, und von P_2 aus die Entfernungen 235,68 m, 235,77 m, 235,71 m gemessen.

Der Mittelwert, das arithmetische Mittel, ist dann

$$l = \frac{1}{6} (235,73 + 235,69 + 235,76 + 235,68 + 235,77 + 235,71) \text{ m} = 235,723 \text{ m}.$$

Bezeichnen wir die Anzahl der Messungen mit n , die einzelnen Meßwerte mit $x_1, x_2, x_3, \dots, x_n$, so ist das arithmetische Mittel $\bar{x}$ gegeben durch

$$\bar{x} = \frac{1}{n} (x_1 + x_2 + x_3 + \dots + x_n) \quad (1)$$

oder unter Verwendung des Summenzeichens durch

$$\bar{x} = \frac{1}{n} \sum_{r=1}^n x_r. \quad (1')$$

Die unmittelbare Anwendung der Formel (1) führt, vor allem wenn die Zahl der Messungen groß wird, zu großen Zahlen. Wir vereinfachen die Rechnung, indem wir von einem angenäherten Mittelwert $\bar{x}_a$, der abgeschätzt wird, die Abweichungen bestimmen und mit dem Mittelwert der Abweichungen das arithmetische Mittel $\bar{x}$ berechnen. Dabei wählt man $\bar{x}_a$ so, daß sich möglichst einfache Zahlen ergeben. Es ist dann

$$\bar{x} = \bar{x}_a + \frac{1}{n} [(x_1 - \bar{x}_a) + (x_2 - \bar{x}_a) + \cdots + (x_n - \bar{x}_a)]$$

oder

$$\bar{x} = \bar{x}_a + \frac{1}{n} \sum_{r=1}^n (x_r - \bar{x}_a). \quad (2)$$

Die Berechnung führen wir stets in einem Rechenschema durch.

Beispiel 2:

Mit einer Feinmeßschraube (Trommelteilung $\frac{1}{100}$ mm) wird der Durchmesser eines Werkstückes siebenmal gemessen.

Messung Nr.	1	2	3	4	5	6	7
Meßgröße in mm	12,357	12,348	12,363	12,341	12,360	12,352	12,368

Es ist der Mittelwert zu bestimmen.

Messung Nr.	Meßwert x_r in mm	Angenäherter Mittelwert $\bar{x}_a$	$(x_r - \bar{x}_a)$ in μ
1	12,357	12,350	+ 7
2	12,348		- 2
3	12,363		+ 13
4	12,341		- 9
5	12,360		+ 10
6	12,352		+ 2
7	12,368		+ 18
Summen			+ 50 - 11 = + 39

$$\bar{x} = \left(12,350 + \frac{1}{7} \cdot \frac{39}{1000} \right) \text{ mm}$$

$$\bar{x} = 12,3556 \text{ mm.}$$

Es genügt nun nicht, die Größe eines Meßobjektes nur durch den Mittelwert aller Einzelmessungen anzugeben. Wir haben ihn zwar als den „richtigen“ Wert bezeichnet. Er kann aber noch einen Fehler enthalten. Die Angabe der Größe nur durch den Mittelwert würde eine Genauigkeit vortäuschen, die nicht vorhanden ist. Wir bestimmen deshalb den **durchschnittlichen Fehler** $\Delta\bar{x}$, das heißt das arithmetische Mittel aus den Absolutbeträgen der Abweichungen vom Mittelwert. Den durchschnittlichen Fehler bezeichnen wir auch als absoluten Fehler. Er gibt an, wie groß die mögliche Abweichung der wirklichen Größe des Meßgegenstandes von dem Mittelwert ist. Dabei muß beachtet werden, daß $\Delta\bar{x}$ positiv oder negativ sein kann. Die Größe des Meßgegenstandes ist dann

$$x = \bar{x} + \Delta\bar{x}. \quad (3)$$

Durch $\bar{x} + \Delta\bar{x}$ ist ein Intervall bestimmt, von dem wir annehmen, daß in ihm der wirkliche Wert der Meßgröße liegt.

In Beispiel 2 beträgt der durchschnittliche Fehler

$$\Delta\bar{x} = \pm \frac{1}{7} (1,4 + 7,6 + 7,4 + 14,6 + 4,4 + 3,6 + 12,4) \mu = \pm 7,3 \mu.$$

Die Größe des Werkstückes muß also mit $x = (12,3556 \pm 0,0073)$ mm angegeben werden.

3. Der relative Fehler

Sind nur die absoluten Fehler der Größen verschiedener Meßgegenstände gegeben, so läßt sich die Genauigkeit der Größenangaben sehr schwer vergleichen. Einen Vergleich ermöglicht das Verhältnis zwischen dem absoluten Fehler und dem Mittelwert der Einzelmessungen (zwischen dem durchschnittlichen Fehler und dem als richtig geltenden Wert). Wir bezeichnen dieses Verhältnis als den **relativen Fehler**. Es ist

$$\sigma = \frac{\Delta\bar{x}}{\bar{x}}. \quad (4)$$

Der relative Fehler beträgt in Beispiel 2

$$\frac{\Delta\bar{x}}{\bar{x}} = \frac{0,0073 \text{ mm}}{12,3556 \text{ mm}} = 0,00059.$$

In der Praxis gibt man häufig den relativen Fehler auch in Prozenten an. Der prozentuale Fehler ist

$$\sigma\% = \frac{\Delta\bar{x}}{\bar{x}} \cdot 100\%. \quad (5)$$

In Beispiel 2 ist der prozentuale Fehler 0,059%.

Aufgaben

1. Mit einer Schieblehre wird ein Werkstück fünfmal gemessen. Dabei ergeben sich die folgenden Meßwerte: 150,73; 150,84; 150,75; 150,64; 150,78 mm.
Es sind der Mittelwert und der durchschnittliche Fehler zu bestimmen.
2. Ein Winkel wird bei einer Vermessung durch sechs Messungen bestimmt. Es ergeben sich die folgenden Werte: $77^{\circ} 48' 43''$, $77^{\circ} 49' 08''$, $77^{\circ} 42' 00''$, $77^{\circ} 49' 38''$, $77^{\circ} 48' 32''$, $77^{\circ} 49' 54''$.
Es sind der Mittelwert und der durchschnittliche Fehler zu bestimmen.
3. Eine Strecke wird mit einem Meßband dreimal gemessen. Dabei werden die folgenden Längen festgestellt: 15,87; 15,72; 15,94 m.
Die mittlere Länge und der mittlere Fehler sind zu bestimmen.
4. Die relativen Fehler sind zu bestimmen
a) in Aufgabe 1, b) in Aufgabe 2, c) in Aufgabe 3.
5. Eine gerade Eisenbahnschiene mit einer Solllänge von 32 m wird mit Meßplatten und einem Gliedemaßstab fünfmal gemessen. Es ergeben sich dabei die folgenden Meßwerte: 31,93; 32,15; 32,07; 31,98; 31,91 m.
Der Mittelwert, der absolute und der relative durchschnittliche Fehler der Meßwerte sowie die Abweichung von der Solllänge sind zu bestimmen.
6. Ein volkseigener Betrieb gewährleistet für Prüfmaschinen zur Untersuchung von Stahl, Eisen und anderen Metallen auf Festigkeit die folgende Genauigkeit:
 1. im ersten Zehntel des Meßbereiches auf ± 1 Teilungsintervall,
 2. darüber auf $\pm 1\%$ des Sollwertes.
 Für jede Prüfmaschine wird vom Werkabnehmer ein Abnahmeprotokoll über die Anzeigegenauigkeit mit 30 Messungen in Sätzen zu je drei Messungen angefertigt. Einem solchen Protokoll sind die folgenden Werte entnommen:

Sollwert	Istwert	Istwert	Istwert	Mittelwert	Abweichung in % des Sollwertes
	2 kp als Einheit				
245,7	246	245	245	245,3	- 0,15
478,3	479	479	479	479,0	+ 0,15
716,7	717	718	718		
958	959	958	958		
1192,7	1192	1193	1193		
1435	1439	1438	1439		
1670	1670	1670	1671		
1910	1910	1910	1910		
2152	2153	2152	2154		
2387	2390	2390	2390		

- Berechnen Sie
- a) die fehlenden Mittelwerte,
 - b) die relativen Abweichungen der Mittelwerte vom zugehörigen Sollwert auf 2 Dezimalen genau,
 - c) den durchschnittlichen relativen Fehler!

7. Ein anderes Abnahmeprotokoll enthält die folgenden Werte:

Sollwert	Istwert	Istwert 6 kp als Einheit	Istwert	Mittelwert	Abweichung in % des Sollwertes
479	479	479	479		
958,7	959	958	959		
1443,3	1443	1444	1444		
1931,7	1937	1936	1936		
2421	2423	2423	2423		
2902	2907	2909	2908		
3392	3395	3394	3395		
3878	3880	3880	3880		
4380	4380	4380	4380		
4850	4858	4858	4859		

Berechnen Sie a) die Mittelwerte in den einzelnen Zeilen,

b) die relativen Abweichungen der Mittelwerte vom Sollwert auf 2 Dezimalen genau,

c) den durchschnittlichen relativen Fehler!

8. Die Dichte einer Flüssigkeit wurde zehnmal nach der gleichen Methode bestimmt. Es ergaben sich die folgenden Werte: 0,9345; 0,9348; 0,9352; 0,9347; 0,9347; 0,9348; 0,9348; 0,9348; 0,9350; 0,9346 g/cm³.

Bestimmen Sie den Mittelwert und den durchschnittlichen Fehler!

4. Der Fehler eines Rechenergebnisses

In der praktischen Mathematik wird, da Meßwerte stets Fehler aufweisen, immer nur mit soviel geltenden Ziffern wie notwendig gerechnet. Dazu ist erforderlich, daß die Genauigkeit des Meßwertes in irgendeiner Weise gekennzeichnet wird und daß entbehrliche Stellen gerundet werden.

a) Rundungs- und Kürzungsregeln

Die Rundungs- und Kürzungsregeln von Zahlen enthält das Normblatt DIN 1333.

1. Rundungsregeln

Die Rundungsregeln beziehen sich auf Zahlen, die in dezimaler Form gegeben sind und die an einer bestimmten Stelle abgebrochen werden sollen.

Abrunden heißt, daß die letzte Stelle, die noch angegeben werden soll, unverändert bleibt. Es wird abgerundet, wenn eine 0, 1, 2, 3 oder 4 folgt.

Beispiel: $8,2134 \approx 8,213 \approx 8,21 \approx 8,2 \approx 8$.

Aufrunden heißt, daß die letzte Stelle, die noch angegeben werden soll, um 1 erhöht wird. Es wird aufgerundet, wenn eine 6, 7, 8 oder 9 folgt.

Beispiel: $3,4579 \approx 3,458 \approx 3,46 \approx 3,5$.

Ist bekannt, daß eine 5 durch Rundung entstanden ist, so wird aufgerundet, wenn die 5 durch Abrundung entstanden ist; ist sie aber durch Aufrunden entstanden, so wird abgerundet.

Beispiele: $6,1852 \approx 6,185 \approx 6,19$
 $6,3146 \approx 6,315 \approx 6,31.$

In allen anderen Fällen wird eine 5 in der letzten Stelle so gerundet, daß die letzte gewünschte Ziffer eine gerade Zahl ist.

Beispiele: $\frac{1}{16} = 0,0625 \approx 0,062$
 $\frac{15}{4} = 3,75 \approx 3,8.$

2. Kürzungsregeln

In der Bruchrechnung verstehen wir unter Kürzen des Bruches die Division des Zählers und des Nenners durch die gleiche Zahl, also eine Formänderung des Bruches, bei der sein Wert nicht geändert wird. In der Fehlerrechnung hat der Begriff **Kürzung** einer Zahl eine andere Bedeutung. Alle Größen, mit denen man in der Fehlerrechnung arbeitet, sind Ergebnisse von Messungen, die mit Fehlern behaftet sind. Um nicht eine Genauigkeit vorzutäuschen, die gar nicht vorhanden ist, gibt man bei Berechnungen die Ausgangswerte nur so genau an, daß die Unsicherheit in der letzten Stelle liegt; alle folgenden Stellen werden unter Berücksichtigung der Rundungsregeln gestrichen. Damit die Größe der Unsicherheit erkennbar ist, bestehen für die Schreibweise gekürzter Werte besondere Regeln.

1. Ist die Unsicherheit in der letzten Stelle höchstens $\pm 0,5$, so wird die Ziffer in gewöhnlichem Schriftgrad gedruckt oder geschrieben.

2. Ist die Unsicherheit in der letzten angegebenen Stelle größer als $\pm 0,5$, so werden die unsicheren Ziffern in Indexstellung oder in kleinerem Schriftgrad angegeben. Steht die letzte in gewöhnlichem Schriftgrad anzugebende Ziffer in der Stelle der Zehner oder davor, so spaltet man Zehnerpotenzen ab.

Beispiel:

Die Masse eines Elektrons ist mit $9,107 \cdot 10^{-28}$ g anzugeben, da die Unsicherheit $0,007 \cdot 10^{-28}$ g beträgt.

In besonderen Fällen kennzeichnet man jedoch die Unsicherheit durch einen Zahlenwert (z. B. bei Toleranzen).

Beispiel: $3,726 \text{ mm} \pm 0,018 \text{ mm} = 3,726 \text{ mm} \pm 18 \mu = (3,726 \pm 0,018) \text{ mm}$
 $= 3,726 \text{ mm} \cdot (1 \pm 0,005) = 3,726 \text{ mm} \cdot (1 \pm 0,5\%).$

Bei einer oberflächlichen Betrachtung scheinen Rundung und Kürzung miteinander übereinzustimmen. Der Unterschied zwischen beiden Begriffen ist aber an den folgenden Beispielen leicht zu erkennen.

Die Zahl π kann auf eine beliebige Anzahl von Stellen genau berechnet werden. Für praktische Rechnungen verwendet man, je nach der geforderten Genauigkeit, nur eine bestimmte Anzahl von Stellen, z. B. 3,14 oder 3,1416. Diese Werte sind auf 2 bzw. 4 Stellen hinter dem Komma gerundet.

Die Lichtgeschwindigkeit im Vakuum ist $c = (2,99790 \pm 0,00006) \cdot 10^{10} \text{ cm} \cdot \text{s}^{-1}$. Für praktische Berechnungen genügt es aber, die Lichtgeschwindigkeit gekürzt mit $2,997_{90} \cdot 10^{10} \text{ cm} \cdot \text{s}^{-1}$ anzugeben. Häufig wird dieser Wert dann noch auf $3 \cdot 10^{10} \text{ cm} \cdot \text{s}^{-1}$ gerundet.

Das Plancksche Wirkungsquantum wurde zu $h = (6,6234 \pm 0,0011) \cdot 10^{-27} \text{ erg} \cdot \text{s}$ gemessen. Dieser Wert kann auf $6,62_3 \cdot 10^{-27} \text{ erg} \cdot \text{s}$ gekürzt werden.

Beim Kürzen wird also ein angegebener Fehler nicht berücksichtigt.

b) Der Fehler einer Funktion

Im Lehrbuch der Mathematik für das 11. Schuljahr (S. 47 bis 49) hatten wir den Fehler berechnet, der sich für den Rauminhalt eines Würfels ergibt, wenn die Kantenlänge des Würfels mit einem bestimmten Fehler behaftet ist. Diesen Gedankengang greifen wir auf, um ihn zu verallgemeinern.

Es sei die Funktion $y = f(x)$ gegeben, wobei x eine Meßgröße mit dem Fehler Δx ist. Der Funktionswert y besitzt dann einen Fehler Δy . Es gilt

$$y + \Delta y = f(x + \Delta x).$$

Durch Subtraktion von $y = f(x)$ erhalten wir den Fehler der Funktion

$$\Delta y = f(x + \Delta x) - f(x).$$

Durch Erweitern der rechten Seite dieser Gleichung mit Δx erhalten wir

$$\Delta y = \frac{f(x + \Delta x) - f(x)}{\Delta x} \cdot \Delta x.$$

Ist Δx hinreichend klein, so kann man den Differenzenquotienten durch den Differentialquotienten ersetzen, und es ergibt sich

$$\Delta y \approx f'(x) \cdot \Delta x. \quad (6)$$

Der relative Fehler ist, wie auch schon gezeigt wurde,

$$\frac{\Delta y}{y} \approx \frac{f'(x)}{f(x)} \cdot \Delta x. \quad (6')$$

Beispiel:

Ein Winkel ist mit $\alpha = 67,37^\circ \pm 0,05^\circ$ gemessen. Wie groß ist der Fehler von $\sin \alpha$?

$$y = \sin(67,37 \pm 0,05)^\circ$$

$$\Delta y \approx \cos 67,37^\circ \cdot \Delta \alpha$$

$$\approx 0,3848 \cdot (\pm 0,0008727) \approx \pm 0,00034.$$

Bei der numerischen Rechnung ist α in das Bogenmaß umzuwandeln. Der Wert $\sin 67,37^\circ = 0,9230$ besitzt also den Fehler $\pm 0,00034$. Hiervon kann man sich unmittelbar überzeugen, wenn man die beiden Werte $\sin 67,32^\circ$ und $\sin 67,42^\circ$ der Funktionstafel entnimmt.

Der relative Fehler beträgt

$$\frac{\Delta y}{y} = \pm \frac{0,00034}{0,9230} = \pm 0,0004, \text{ das heißt } \pm 0,04\%.$$

c) Der Fehler einer Summe

Eine Summe sei bestimmt durch die Gleichung

$$y = a_1 + a_2 + a_3 + \cdots + a_n,$$

wobei die Summanden $a_1, a_2, \dots, a_n$ Meßgrößen mit den Fehlern $\Delta a_1, \Delta a_2, \dots, \Delta a_n$ seien. Die Addition führt auf einen Wert y , der den Fehler Δy besitzt. Da wir nicht wissen, welche der Fehler Δa_i positiv und welche negativ sind, können wir nichts darüber aussagen, ob der Fehler eines Summanden den des Ergebnisses vergrößert oder verkleinert, und auch nichts darüber, welchen Einfluß ein Subtrahend mit seinem Fehler hat; wir können nur feststellen, welchen Wert die Summe unter Berücksichtigung der möglichen Einzelfehler höchstens besitzt. Es ist

$$y + \Delta y = a_1 + \Delta a_1 + a_2 + \Delta a_2 + \cdots + a_n + \Delta a_n.$$

Durch Subtraktion von $y = a_1 + a_2 + a_3 + \cdots + a_n$ erhalten wir

$$\Delta y = \Delta a_1 + \Delta a_2 + \Delta a_3 + \cdots + \Delta a_n.$$

Einige Glieder können in dieser Summe negativ sein, sie können also den Wert des Fehlers verkleinern. Da wir jedoch den größtmöglichen (maximalen) Fehler bestimmen wollen, müssen wir zu den Absolutbeträgen übergehen. Wir erhalten also zunächst

$$|\Delta y| \leq |\Delta a_1| + |\Delta a_2| + |\Delta a_3| + \cdots + |\Delta a_n|.$$

Der größtmögliche (maximale) Fehler ist daher

$$|\Delta y| = |\Delta a_1| + |\Delta a_2| + |\Delta a_3| + \cdots + |\Delta a_n|. \quad (7)$$

Der Wert des größten möglichen Fehlers einer Summe ist gleich der Summe aus den Absolutbeträgen der möglichen Fehler aller Summanden.

Dividieren wir den absoluten Fehler der Summe durch deren Wert, so erhalten wir den relativen Fehler. Dabei sind jedoch zwei Fälle zu unterscheiden.

Besteht die Summe nur aus positiven Gliedern, so gibt der relative Fehler einen zuverlässigen Einblick in die Größe des Fehlers.

Besitzt die Summe negative Glieder und ist ihr Gesamtbetrag klein gegenüber den einzelnen Gliedern, so kann der relative Fehler außergewöhnlich große Werte annehmen.

Beispiel 1:

Zwei Strecken sind mit $(62,4 \pm 0,5)$ mm und mit $(10,7 \pm 0,5)$ mm gemessen. Wie groß ist der relative Fehler der Differenz?

Es ist

$$a_1 - a_2 = 51,7 \text{ mm und } |\Delta a_1| + |\Delta a_2| = 1 \text{ mm,}$$

$$\frac{\Delta y}{y} = \pm \frac{1 \text{ mm}}{51,7 \text{ mm}} = \pm 0,019,$$

das heißt, der relative Fehler beträgt $\pm 1,9\%$.

Beispiel 2:

Zwei Strecken sind mit $(15,7 \pm 0,5)$ mm und mit $(14,5 \pm 0,5)$ mm gemessen. Welchen relativen Fehler besitzt die Differenz?

Es ist

$$a_1 - a_2 = 1,2 \text{ mm und } |\Delta a_1| + |\Delta a_2| = 1 \text{ mm,}$$

$$\frac{\Delta y}{y} = \pm \frac{1 \text{ mm}}{1,2 \text{ mm}} = \pm 0,83,$$

das heißt, der relative Fehler beträgt $\pm 83\%$.

d) Der Fehler eines Produktes

Es sei $y = u \cdot v$ das Produkt aus den Meßwerten u und v mit den voneinander unabhängigen Meßfehlern Δu und Δv . Infolgedessen ist das Ergebnis fehlerhaft. Es gilt

$$y + \Delta y = (u + \Delta u)(v + \Delta v).$$

Hieraus ergibt sich

$$y + \Delta y = u \cdot v + v \cdot \Delta u + u \cdot \Delta v + \Delta u \cdot \Delta v$$

und damit nach Subtraktion von $y = u \cdot v$

$$\Delta y = v \cdot \Delta u + u \cdot \Delta v + \Delta u \cdot \Delta v.$$

Wir wollen wieder den maximalen Fehler bestimmen und gehen deshalb zu den Absolutbeträgen über (vgl. Abschnitt c). Damit erhalten wir

$$|\Delta y| \leq |v \cdot \Delta u| + |u \cdot \Delta v| + |\Delta u \cdot \Delta v|.$$

Der maximale Fehler ist daher

$$|\Delta y| = |v \cdot \Delta u| + |u \cdot \Delta v| + |\Delta v \cdot \Delta u|.$$

In vielen Fällen ist das Produkt $|\Delta u \cdot \Delta v|$ so klein, daß es gegen die übrigen Glieder der Summe vernachlässigt werden kann. In den Aufgaben zu diesem Abschnitt ist dies stets der Fall. Wir setzen deshalb

$$|\Delta y| = |v \cdot \Delta u| + |u \cdot \Delta v|. \quad (8)$$

Bei jedem praktischen Arbeiten muß jedoch die Größenordnung des Produktes $|\Delta u \cdot \Delta v|$ abgeschätzt werden.

Dividieren wir die Formel (8) durch $|y| = |u \cdot v|$, so ergibt sich der relative Fehler zu

$$\left| \frac{\Delta y}{y} \right| = \left| \frac{\Delta u}{u} \right| + \left| \frac{\Delta v}{v} \right|. \quad (8')$$

Der Betrag des relativen Fehlers eines Produktes $y = u_1 \cdot u_2 \cdot u_3 \cdot \dots \cdot u_n$, dessen Faktoren $u_1, u_2, u_3, \dots, u_n$ die Fehler $\Delta u_1, \Delta u_2, \Delta u_3, \dots, \Delta u_n$ haben, beträgt

$$\left| \frac{\Delta y}{y} \right| = \left| \frac{\Delta u_1}{u_1} \right| + \left| \frac{\Delta u_2}{u_2} \right| + \left| \frac{\Delta u_3}{u_3} \right| + \dots + \left| \frac{\Delta u_n}{u_n} \right|. \quad (8'')$$

Diese Formel läßt sich rekursiv aus der Formel (8') herleiten, wenn man jeweils zwei Faktoren zusammenfaßt.

Der relative Fehler eines Produktes ist gleich der Summe aus den relativen Fehlern der Faktoren.

e) Der Fehler eines Quotienten

Gegeben sei der Quotient $y = \frac{u}{v}$ zweier Meßwerte u und v , die mit den Fehlern Δu und Δv behaftet sind.

Dann ist

$$y + \Delta y = \frac{u + \Delta u}{v + \Delta v}.$$

Es folgt

$$\Delta y = \frac{u + \Delta u}{v + \Delta v} - \frac{u}{v}$$

oder

$$\Delta y = \frac{v \cdot \Delta u - u \cdot \Delta v}{v(v + \Delta v)}.$$

Nun interessiert uns nur der maximale Fehler Δy . Deshalb untersuchen wir, unter welchen Bedingungen der Bruch einen maximalen Wert annimmt.

Werden bei einem Bruch Zähler und Nenner unabhängig voneinander verändert, so kann man über die Wertänderung des Bruches keine Aussage machen, wenn

Zähler und Nenner entweder beide vergrößert oder beide verkleinert werden. Wird dagegen der Zähler vergrößert und der Nenner verkleinert, so wird der Wert des Bruches sicher vergrößert. Deshalb setzen wir beim Übergang zu den Absolutbeträgen im Zähler das Pluszeichen und im Nenner das Minuszeichen. Damit erhalten wir für den maximalen Fehler

$$|\Delta y| = \frac{|v \cdot \Delta u| + |u \cdot \Delta v|}{|v| \cdot (|v| - |\Delta v|)}.$$

Ziehen wir $|v|$ im Nenner vor die Klammer, so erhalten wir

$$|\Delta y| = \frac{|v \cdot \Delta u| + |u \cdot \Delta v|}{v^2 \left(1 - \left|\frac{\Delta v}{v}\right|\right)}$$

oder

$$|\Delta y| = \frac{|v \cdot \Delta u| + |u \cdot \Delta v|}{v^2} \cdot \frac{1}{\left(1 - \left|\frac{\Delta v}{v}\right|\right)}.$$

Ist $|\Delta v|$ klein gegen $|v|$, so unterscheidet sich der Ausdruck $\frac{1}{1 - \left|\frac{\Delta v}{v}\right|}$ wenig von 1. Wir können ihn entsprechend den Betrachtungen bei der Herleitung des Fehlers eines Produktes vernachlässigen. Der maximale Fehler ist also

$$|\Delta y| = \frac{|v \cdot \Delta u| + |u \cdot \Delta v|}{v^2}. \quad (9)$$

Nach Division durch $|y| = \left|\frac{u}{v}\right|$ erhalten wir

$$\left|\frac{\Delta y}{y}\right| = \frac{|v \cdot \Delta u| + |u \cdot \Delta v|}{v^2 \cdot \left|\frac{u}{v}\right|}.$$

$$= \frac{|v \cdot \Delta u| + |u \cdot \Delta v|}{|u \cdot v|}$$

$$\left|\frac{\Delta y}{y}\right| = \left|\frac{\Delta u}{u}\right| + \left|\frac{\Delta v}{v}\right|. \quad (9')$$

Der relative Fehler eines Quotienten ist gleich der Summe aus den relativen Fehlern des Divisors und des Dividenten.

f) Anwendung in der Trigonometrie

In einem rechtwinkligen Dreieck seien die Hypotenuse c mit dem Meßfehler Δc und der Winkel α mit dem Meßfehler $\Delta \alpha$ bekannt. Es ist der Fehler der Kathete a zu bestimmen, wenn $c = (25,00 \pm 0,05)$ cm und $\alpha = (30 \pm 0,2)^\circ$ sind.

Da $a = c \cdot \sin \alpha$ ist, errechnen wir den Fehler für a unter Verwendung der Formel (8)

$$|\Delta y| = |v \cdot \Delta u| + |u \cdot \Delta v|.$$

Wir setzen $u = c$ und $v = \sin \alpha$. Dann ist $\Delta u = \Delta c$. Für Δv erhalten wir nach Formel (6)

$$\Delta v \approx \cos \alpha \cdot \Delta \alpha.$$

Für Δa ergibt sich dann

$$|\Delta a| = |\sin \alpha \cdot \Delta c| + |c \cdot \cos \alpha \cdot \Delta \alpha|.$$

Für den relativen Fehler erhalten wir

$$\left| \frac{\Delta a}{a} \right| = \left| \frac{\Delta c}{c} \right| + \left| \frac{\cos \alpha \cdot \Delta \alpha}{\sin \alpha} \right| = \left| \frac{\Delta c}{c} \right| + |\operatorname{ctg} \alpha \cdot \Delta \alpha|.$$

Bei der numerischen Berechnung ist zu beachten, daß α in das Bogenmaß umzurechnen ist.

Es ergibt sich

$$\left| \frac{\Delta a}{a} \right| = \left| \frac{0,05}{25} \right| + |1,732 \cdot 0,00349| = |0,002| + |0,0060|$$

$$\frac{\Delta a}{a} = \pm 0,0080.$$

Der relative Fehler beträgt also $\pm 0,80\%$.

Enthält eine Funktion mehrere unabhängige Veränderliche, so bestimmen wir nach der Formel (6) die Fehler der einzelnen Veränderlichen und setzen den absoluten Fehler unter Verwendung der Formeln (7), (8) bzw. (9) zusammen.

Soll der relative Fehler bestimmt werden, so dividieren wir beide Seiten des absoluten Fehlers durch die gegebene Funktion, wobei wir soweit wie möglich vereinfachen.

Den Fehler von Winkelfunktionen beziehen wir auf das Bogenmaß.

Aufgaben

Es sind, wenn in der Aufgabe nichts anderes angegeben ist, die absoluten und die relativen Fehler zu berechnen.

1. Die Seite eines Quadrates ist $a = (52,52 \pm 0,05)$ mm. Berechnen Sie die Fläche und ihren Fehler!
2. Der Durchmesser eines Kreises ist mit $d = 13,43$ cm, $13,37$ cm, $13,39$ cm gemessen. Wie groß ist a) der mittlere Durchmesser, b) der Meßfehler, c) der Inhalt und dessen Fehler?
8. Die Seiten eines Dreiecks werden mit Hilfe eines Stechzirkels auf einem Transversal-Maßstab bestimmt. Dabei findet man die Werte $a = 7,35$ cm, $b = 5,27$ cm, $c = 8,64$ cm. Die Meßwerte der Seiten sind mit $\pm 0,1$ mm unsicher. Wie groß ist der Fehler des Umfanges?
4. In einem Gleichstromkreis wurde die Spannung mit $U = (36,53 \pm 0,005)$ Volt und die Stromstärke mit $I = (0,056 \pm 0,0005)$ Ampere gemessen. Der Widerstand des Stromkreises ist zu bestimmen.
5. Die Seiten eines Rechtecks sind mit $81,5$ mm und $57,5$ mm gemessen. Der Fehler jeder Länge betrage $\pm 0,2$ mm. Wie groß sind der absolute und der relative Fehler des Flächeninhaltes?
6. Die Grundseite eines Dreiecks wurde mit $4,57$ cm und die zugehörige Höhe mit $3,14$ cm bestimmt. Der Fehler der Meßwerte betrage je $\pm 0,2$ mm. Wie groß sind der absolute und der relative Fehler des Flächeninhaltes?
7. An einem Trapez werden die Grundlinien mit $a = 7,45$ cm und $b = 5,79$ cm und die Höhe mit $h = 4,25$ cm gemessen. Der Meßfehler sei $\pm 0,5$ mm. Wie groß ist der Fehler des Flächeninhaltes?
8. Bei einem Litermaß aus Aluminium sind der Durchmesser und die Höhe je dreimal gemessen worden. Es ergaben sich die folgenden Meßwerte:
 $d = 8,55$ cm, $8,57$ cm, $8,56$ cm und $h = 17,43$ cm, $17,36$ cm, $17,35$ cm.
 Bestimmen Sie a) die Mittelwerte von d und h , b) die durchschnittlichen Fehler, c) das Volumen und dessen Fehler, d) den Fehler des Volumens gegenüber dem Sollwert!
 Überlegen Sie, wieviel Stellen von π verwendet werden müssen, damit durch dessen Näherungswert der Fehler des Ergebnisses nur unwesentlich beeinflußt wird!
9. Beim Ausmessen eines Ziegelsteines sind auf je zwei gegenüberliegenden Flächen die folgenden Maße in cm festgestellt worden:

a	b	c
25,35	12,00	6,58
25,32	12,10	6,70
25,28	12,05	6,62
25,48	11,92	6,72
25,40	12,20	6,63
25,28	12,12	6,78

- Bestimmen Sie a) die Mittelwerte der Kantenlängen,
 b) den durchschnittlichen Fehler der Messungen,
 c) den absoluten und den relativen Fehler des Rauminhaltes,
 d) die Abweichung des Volumens vom Sollwert (Kantenlängen 25 cm, 12 cm und 6,5 cm)!

10. Bei einem Blatt DIN A 4 sind durch je drei Messungen die folgenden Werte ermittelt worden:

$a = 210,3 \text{ mm}$, $209,8 \text{ mm}$, $211,1 \text{ mm}$ und $b = 296,5 \text{ mm}$, $296,4 \text{ mm}$, $296,3 \text{ mm}$.

- Bestimmen Sie a) die Mittelwerte der Seiten und deren Fehler,
 b) die durchschnittlichen Fehler gegenüber dem Sollmaß,
 c) die Fläche und deren Fehler,
 d) den Fehler gegenüber der Sollfläche!

11. Es ist im rechtwinkligen Dreieck $b = c \cdot \cos \alpha$. Bestimmen Sie den absoluten und den relativen Fehler der Seite b ! Es sei $c = 203,72 \text{ m} \pm 0,05 \text{ m}$, $\alpha = 37,37^\circ \pm 0,05^\circ$.

12. Es ist im rechtwinkligen Dreieck $\tan \alpha = \frac{a}{b}$. Bestimmen Sie den absoluten und den relativen Fehler von α ! Es sei $a = 25,73 \text{ m} \pm 0,05 \text{ m}$, $b = 18,57 \text{ m} \pm 0,05 \text{ m}$.

13. Der Flächeninhalt eines Dreiecks ist $F = \frac{ab \sin \gamma}{2}$. Bestimmen Sie den Fehler, wenn $a = 20,00 \text{ cm} \pm 0,02 \text{ cm}$, $b = 10,00 \text{ cm} \pm 0,02 \text{ cm}$, $\gamma = 30^\circ \pm 0,2^\circ$ ist!

14. Nach dem Sinussatz ist $a = \frac{b \cdot \sin \alpha}{\sin \beta}$. Bestimmen Sie den Fehler von a , wenn $b = 10,00 \text{ cm} \pm 0,02 \text{ cm}$, $\alpha = 55^\circ \pm 0,2^\circ$, $\beta = 65^\circ \pm 0,2^\circ$ ist!

15. Die Masse eines kleinen Quecksilbertropfens soll in einem Uhrglas mit einer analytischen Waage bestimmt werden. Man erhielt die folgenden Mittelwerte:

Masse des Uhrglases ohne Hg $(6102,37 \pm 0,05) \text{ mg}$

Masse des Uhrglases mit Hg $(6109,21 \pm 0,05) \text{ mg}$.

Wie groß sind die Masse des Quecksilbertropfens und der relative Fehler der Wägung?

16. Die Kapazität eines Plattenkondensators soll berechnet werden. Die Fläche der Platten wurde mit $F = (105,7 \pm 0,13) \text{ cm}^2$ und der Abstand der Platten mit $(0,1 \pm 0,01) \text{ cm}$ bestimmt.

Die Dielektrizitätskonstante ist $\epsilon_0 = 8,859 \cdot 10^{-14} \frac{\text{Coulomb}}{\text{Volt} \cdot \text{cm}}$.

17. Die Kapazität eines Kugelkondensators beträgt $C = 4\pi\epsilon_0 \frac{R \cdot r}{R - r}$.

Als Mittelwerte der Messungen ergaben sich: $R = (5,27 \pm 0,01) \text{ cm}$ und $r = (5,012 \pm 0,002) \text{ cm}$. Es ist der relative Fehler der Kapazität zu bestimmen.

II. Näherungslösungen

5. Praktische Berechnung erforderlicher Funktionswerte

In den Betrachtungen, die wir im folgenden durchführen werden, ist es häufig notwendig, Funktionswerte $f(x)$ von ganzen rationalen Funktionen zu berechnen. Diese Berechnung wird mitunter bereits dann umständlich, wenn $f(x)$ dritten Grades ist. Wir entwickeln daher ein Schema, das die Berechnung mit Hilfe des Rechenstabes gestattet. Dazu gehen wir von der Funktion dritten Grades

$$f(x) = a_3 x^3 + a_2 x^2 + a_1 x + a_0$$

(mit $a_3 \neq 0$) als Beispiel aus. Wir ziehen den Faktor x vor die Klammer und erhalten damit

$$f(x) = [a_3 x^2 + a_2 x + a_1] x + a_0.$$

Durch weitere Abtrennung des Faktors x ergibt sich

$$f(x) = [(a_3 x + a_2) x + a_1] x + a_0.$$

Diese Gleichung ermöglicht es uns, den Funktionswert $f(x)$ in aufeinanderfolgenden Stufen zu berechnen:

1. Den Koeffizienten a_3 multiplizieren wir mit x .
2. Zum Produkt addieren wir den Koeffizienten a_2 .
3. Die Summe multiplizieren wir mit x .
4. Zu diesem Produkt addieren wir a_1 .
5. Diese Summe multiplizieren wir mit x .
6. Zum Produkt addieren wir a_0 .

Diesen Rechengang stellen wir in dem folgenden Schema dar:

a_3	a_2	a_1	a_0
1. $a_3 x$	2. $a_3 x + a_2$		
	3. $(a_3 x + a_2) x$	4. $(a_3 x + a_2) x + a_1$	
		5. $[(a_3 x + a_2) x + a_1] x$	6. $[(a_3 x + a_2) x + a_1] x + a_0$

Verkürzt schreiben wir das Schema folgendermaßen:

a_3	a_2	a_1	a_0
	$a_3 x$	$(a_3 x + a_2) x$	$[(a_3 x + a_2) x + a_1] x$
a_3	$a_3 x + a_2$	$(a_3 x + a_2) x + a_1$	$[(a_3 x + a_2) x + a_1] x + a_0$

Entsprechend kann dieses Schema für ganze rationale Funktionen n -ten Grades aufgestellt werden. Es heißt **Horner'sches Schema**.

Beispiel:

Es soll $f(x) = 2,5x^3 - 3,1x^2 + 1,5x + 0,8$ für $x = 1,5$ berechnet werden. Wir schreiben die Koeffizienten der Funktion in eine Zeile hintereinander:

$$2,5 \quad -3,1 \quad 1,5 \quad 0,8.$$

Nun multiplizieren wir $2,5$ mit $x=1,5$. Wir erhalten $3,75$ und addieren diesen Wert zu $-3,1$. Das Ergebnis $0,65$ multiplizieren wir wieder mit $x=1,5$. Das Produkt fassen wir mit dem nächsten Koeffizienten $1,5$ zusammen usw. Es ergibt sich dann im Schema:

$$\begin{array}{ccccccc}
 2,5 & -3,1 & & 1,5 & & 0,8 & \\
 & \nearrow 3,75 & & \nearrow 0,98 & & \nearrow 3,72 & \\
 \hline
 2,5 & & 0,65 & & 2,48 & & 4,5
 \end{array}
 \quad f(1,5) = 4,5.$$

Im Rechenstab stellen wir $1,5$ fest ein und lesen sämtliche Produkte ab, ohne die Zunge zu versetzen.

Sind in der Funktion ein oder mehrere Koeffizienten gleich Null, so werden die Koeffizienten 0 in der ersten Zeile mitgeschrieben.

Beispiel:

Es ist $f(x) = 5x^4 - 2x^2 - 5x$ für $x = 2$ zu berechnen.

$$f(x) = 5x^4 + 0 \cdot x^3 - 2x^2 - 5x + 0$$

$$\begin{array}{cccccc}
 5 & 0 & -2 & -5 & 0 & \\
 & 10 & 20 & 36 & 62 & \\
 \hline
 5 & 10 & 18 & 31 & 62 &
 \end{array}
 \quad f(2) = 62.$$

6. Das Sekantennäherungsverfahren

Gegeben sei eine Funktion $y = f(x)$. Es soll eine Nullstelle dieser Funktion bestimmt werden.

Zunächst untersuchen wir, ob die Funktion überhaupt eine Nullstelle hat. Dazu stellen wir eine Wertetafel auf; an dieser prüfen wir, ob mindestens zwei Funktionswerte vorhanden sind, die entgegengesetzte Vorzeichen haben. Existieren zwei solche Funktionswerte $f(x_1)$ und $f(x_2)$, so muß noch festgestellt werden, ob die Funktion im Intervall von x_1 bis x_2 stetig ist. Wenn das der Fall ist, so hat sie in diesem Intervall mindestens eine Nullstelle x_0 .

Bei allen folgenden Betrachtungen werden nur solche Funktionen behandelt, die höchstens Lücken oder Pole als Unstetigkeitsstellen aufweisen. Es genügt also, die Untersuchung auf diese Fälle zu beschränken.

Durch die Punkte $[x_1; f(x_1)]$ und $[x_2; f(x_2)]$ legen wir die Sekante (Abb. 1); ihre Gleichung ist

$$\frac{f(x_2) - f(x_1)}{x_2 - x_1} = \frac{\eta - f(x_1)}{\xi - x_1},$$

wenn mit ξ und η die Koordinaten aller Punkte auf der Sekante bezeichnet werden. Für $\eta_1 = 0$ hat die Sekante eine Nullstelle ξ_1 . Es gilt dann

$$\xi_1 = x_1 - \frac{x_2 - x_1}{f(x_2) - f(x_1)} \cdot f(x_1). \quad (1)$$

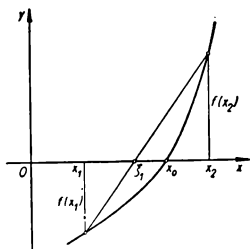


Abb. 1

Da x_0 und ξ_1 Werte im Intervall $x_1 \dots x_2$ sind, liegt ξ_1 näher an x_0 als mindestens einer der beiden Werte x_1 und x_2 .

Durch Fortsetzung dieses Verfahrens kann man die Stelle x_0 mit beliebiger Genauigkeit annähern. Dabei wählt man stets die Ausgangswerte für die weitere Berechnung so, daß die zugehörigen Funktionswerte verschiedene Vorzeichen haben.

Beispiel:

Gegeben sei die Funktion $y = f(x) = x^3 + x - 5 = 0$.

x	0	+1	+2	+3
y	-5	-3	+5	+25

Aus der Wertetafel erkennen wir, daß $f(1) = -3$ und $f(2) = +5$ verschiedene Vorzeichen haben. Da es sich bei unserem Beispiel um eine ganze rationale Funktion handelt, können Lücken und Pole nicht auftreten. Daraus folgt, daß zwischen $x_1 = 1$ und $x_2 = 2$ eine Nullstelle liegt. Für ξ_1 ergibt sich

$$\xi_1 = 1 - \frac{2-1}{5-(-3)} \cdot (-3) = 1 + \frac{3}{8} = 1,375.$$

Den Wert $\xi_1 = 1,375$ runden wir auf $\xi_1 \approx 1,4$, da ξ_1 ohnehin nur ein Näherungswert ist und die Verwendung des Wertes $\xi_1 = 1,375$ die Fortsetzung der Rechnung erschwert.

Für $f(1,4)$ ergibt sich

$$f(1,4) \approx -0,86.$$

Da $f(1,4)$ negativ ist, wählen wir den positiven Wert $f(x_2) = 5$ für die Berechnung des nächsten Näherungswertes. Wir bezeichnen ihn mit ξ_2 , und es gilt

$$\xi_2 = \xi_1 - \frac{x_2 - \xi_1}{f(x_2) - f(\xi_1)} \cdot f(\xi_1).$$

Diese Gleichung wurde aus der Gleichung für ξ_1 gewonnen, indem dort für x_1 der bessere Näherungswert ξ_1 eingesetzt wurde. Für ξ_2 erhält man

$$\xi_2 = 1,4 - \frac{2 - 1,4}{5 - (-0,86)} \cdot (-0,86) = 1,4 + \frac{0,516}{5,86} \approx 1,5.$$

Für $f(1,5)$ ergibt sich

$$f(1,5) = -0,125.$$

Führt man eine weitere Näherung durch, so erhält man $\xi_3 = 1,512$ und $f(1,512) \approx -0,031$. Es gilt also

$$x_0 \approx 1,512.$$

Das eben behandelte Verfahren heißt **Sekantennäherungsverfahren** oder auch **regula falsi**. Diese Bezeichnung (mittelalterliches Latein) bedeutet soviel wie „Regel vom Falschen (ausgehend)“.

7. Das Tangentennäherungsverfahren

Bei einem weiteren Verfahren zur näherungsweise Bestimmung von Nullstellen einer Funktion wird an Stelle der Sekante die Tangente verwendet.

Es sei $f(x)$ eine Funktion, von der bekannt ist, daß sie im Intervall $x_1 \leq x \leq x_2$ stetig und mindestens zweimal differenzierbar ist (vgl. die einleitenden Absätze zum Abschnitt 6) und daß sie in diesem Intervall eine Nullstelle x_0 besitzt. Ferner sei $f'(x)$ und $f''(x)$ im ganzen Intervall verschieden von Null, das heißt, in diesem Intervall existieren weder Extremwerte noch Wendepunkte. Es haben also $f(x_1)$ und $f(x_2)$ verschiedene Vorzeichen.

Die Gleichung der Tangente an die Kurve der Funktion $f(x)$ im Punkte $[x_1; f(x_1)]$ ist

$$\frac{\eta - f(x_1)}{\xi - x_1} = f'(x_1),$$

wobei mit ξ und η die Koordinaten aller Punkte auf der Tangente bezeichnet sind. Für den Schnittpunkt der Tangente mit der x -Achse ergibt sich:

$$\eta_1 = 0; \quad \xi_1 = x_1 - \frac{f(x_1)}{f'(x_1)}.$$

Wir müssen nun untersuchen, ob ξ_1 einen besseren Näherungswert für x_0 darstellt als x_1 .

Wenn $f(x_1)$ und $f''(x_1)$ beide negativ sind, so steigt die Kurve von x_1 bis x_2 , und sie ist konkav von unten. Die Tangente liegt also oberhalb der Kurve und schneidet demnach die x -Achse zwischen x_1 und x_0 (Abb. 2). Zum entsprechenden Ergebnis gelangen wir, wenn $f(x_1)$ und $f''(x_1)$ beide positiv sind (Abb. 3). In diesen Fällen

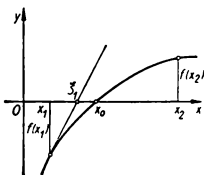


Abb. 2

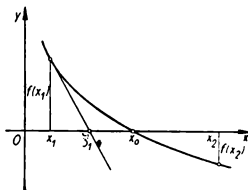


Abb. 3

stellt ξ_1 eine bessere Näherung für x_0 dar als x_1 . Sind dagegen die Vorzeichen von $f(x_1)$ und $f''(x_1)$ verschieden (Abb. 4 und 5), so liegt der Schnittpunkt nicht zwischen x_1 und x_0 . Über seine Brauchbarkeit als Näherungswert kann man daher keine Aussage machen.

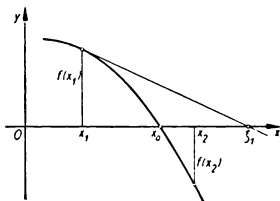


Abb. 4

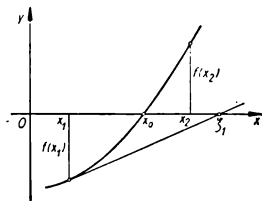


Abb. 5

Eine entsprechende Betrachtung kann man für den Näherungswert x_2 anstellen.

Da $f(x_1)$ und $f(x_2)$ verschiedene Vorzeichen haben, stimmt das Vorzeichen eines und nur eines dieser beiden Werte mit dem Vorzeichen von $f''(x)$ überein (das sich nach der Voraussetzung im ganzen Intervall nicht ändert). Daher ist nur einer der beiden Werte $f(x_1)$ und $f(x_2)$ für die Berechnung von ξ_1 sicher brauchbar.

Es ist

$$\xi_1 = x_1 - \frac{f(x_1)}{f''(x_1)} \quad (2)$$

ein besserer Näherungswert für x_0 als der Wert x_1 , wenn

$$f(x_1) \cdot f''(x_1) > 0$$

ist.

Wir haben diese Bedingung aus der geometrischen Veranschaulichung gewonnen. Auf einen strengen analytischen Beweis soll wegen seines Umfanges verzichtet werden.

Durch Fortsetzung dieses Verfahrens kann man die Stelle x_0 mit beliebiger Genauigkeit annähern.

Beispiel:

Wir betrachten wieder die Funktion

$$f(x) = x^3 + x - 5.$$

Zwei Näherungswerte sind $x_1 = 1$ und $x_2 = 2$. Es ist

$$f'(x) = 3x^2 + 1$$

und

$$f''(x) = 6x.$$

Im Intervall $1 \leq x \leq 2$ sind beide Ableitungen positiv, also verschieden von Null. Wegen $f'(x) > 0$ wählen wir $f(2) = 5$ als ersten Näherungswert. Es ist dann

$$\xi_1 = 2 - \frac{5}{13} \approx 1,6$$

und

$$f(1,6) \approx 0,7.$$

Verwenden wir für die Berechnung des nächsten Näherungswertes diesen Wert, so erhalten wir

$$\xi_2 = 1,6 - \frac{0,7}{8,7} \approx 1,52$$

und

$$f(1,52) \approx 0,03.$$

Der dritte Näherungswert ergibt sich als

$$\xi_3 = 1,52 - \frac{0,03}{7,93} \approx 1,516$$

mit

$$f(1,516) \approx -0,00016.$$

Das eben behandelte Verfahren, das **Tangentennäherungsverfahren**, wurde von *Isaak Newton* entwickelt und wird nach ihm auch als **Newton'sches Näherungsverfahren** bezeichnet.

Das Newton'sche Näherungsverfahren führt schneller zu einem guten Näherungswert als die *regula falsi*. Es hat aber gegenüber der *regula falsi* den Nachteil, daß für seine Anwendbarkeit mehr Einschränkungen gemacht werden mußten (vgl. dazu die einleitenden Absätze zu beiden Verfahren).

Da die Ableitungen bei ganzen rationalen Funktionen leicht zu bestimmen sind, wird man in diesen Fällen in der Regel das Newton'sche Näherungsverfahren anwenden. Liegt jedoch ein Extremwert oder ein Wendepunkt im betrachteten

Intervall, so ist es zweckmäßig, mit Hilfe der regula falsi die ersten Näherungswerte zu bestimmen. Gelingt es, auf diese Weise ein Intervall abzugrenzen, in dem zwar eine Nullstelle liegt, aber keine Extremwerte oder Wendepunkte mehr auftreten, so setzt man die Berechnung mit Hilfe des Newtonschen Näherungsverfahrens fort. Ist eine gebrochene rationale oder eine nichtrationale Funktion gegeben, so ist das Newtonsche Näherungsverfahren nur dann zweckmäßig, wenn die Ableitungen leicht zu bestimmen sind. In der Regel wird man in diesem Fall also die regula falsi anwenden. Dabei ist zu beachten, daß man bei gebrochenen Funktionen zur Bestimmung von Nullstellen zunächst nur die Zählerfunktion untersucht und dann erst prüft, ob die Nullstelle der Zählerfunktion auch Nullstelle der Nennerfunktion ist.

Aufgaben

1. Bestimmen Sie eine Wurzel der folgenden Gleichungen

1. nach der regula falsi,
2. nach dem Newtonschen Näherungsverfahren!

Anmerkung: Berechnen Sie alle Funktionswerte mit dem Rechenstab unter Verwendung des Hornerschen Schemas!

- | | |
|-------------------------------------|---------------------------------------|
| a) $x^3 - 5x - 15 = 0$ | b) $x^3 + 3x - 7 = 0$ |
| c) $x^3 + x^2 - 5 = 0$ | d) $2x^3 - 4x - 3 = 0$ |
| e) $5x^3 - 6x - 7 = 0$ | f) $2,5x^3 - 3,1x^2 + 1,5x + 0,8 = 0$ |
| g) $x^4 - 2x^3 + 3x^2 + x - 10 = 0$ | h) $x^5 + 6x^3 + 5x^2 + 10x - 18 = 0$ |
| i) $x^5 - 4x^4 + 3x^3 - x + 9 = 0$ | k) $x^5 - 3x^3 - 2x^2 - 5 = 0$ |

2. Bestimmen Sie eine Wurzel der folgenden Gleichungen

1. nach der regula falsi,
2. nach dem Newtonschen Näherungsverfahren!

- | | |
|-----------------------------------|----------------------------------|
| a) $\cos x - x = 0$ | b) $2 \sin x - x = 0$ |
| c) $\operatorname{ctg} x - x = 0$ | d) $\operatorname{tg} x - x = 0$ |
| e) $7 \sin x - x + 9 = 0$ | f) $2x + \frac{1}{x} - 4 = 0$ |

Anleitung: Die zeichnerische Lösung der Gleichung ergibt einen ersten Näherungswert für eine reelle Wurzel. Vergleichen Sie zur zeichnerischen Lösung der Gleichungen den Abschnitt Goniometrie, Lehrbuch der Mathematik für das 10. Schuljahr!

3. Wie tief sinkt eine Eisenkugel von 10 cm Radius in Quecksilber ein? (Die Wichte von Eisen ist $7,8 \text{ p/cm}^3$, von Quecksilber $13,55 \text{ p/cm}^3$).
4. Einer Halbkugel mit dem Radius $r = 10 \text{ cm}$ soll ein Zylinder einbeschrieben werden, dessen Rauminhalt $\frac{1}{3}$ des Halbkugelvolumens beträgt. Wie groß sind der Radius und die Höhe des Zylinders?
5. Einer Halbkugel soll ein Kegel, dessen Spitze im Kugelmittelpunkt liegt, einbeschrieben werden. Der Radius der Kugel beträgt $r = 12 \text{ cm}$. Wie groß sind der Radius und die Höhe des Kegels, wenn sein Volumen 500 cm^3 betragen soll?

B. Potenzreihen

I. Darstellung einer Funktion durch Potenzreihen

Von allen uns bisher bekannten Funktionen sind die rationalen Funktionen dadurch besonders gekennzeichnet, daß ihre numerische Berechnung verhältnismäßig einfach ist. Der Funktionswert y läßt sich nämlich bei gegebenem x in endlich vielen Schritten durch Anwendung rationaler Rechenoperationen, das heißt durch Addition, Subtraktion, Multiplikation und Division, auf die Variable x berechnen. Bei jeder anderen Funktion ist die Berechnung der Funktionswerte in endlich vielen Schritten im allgemeinen nicht möglich. So braucht man zum Beispiel zur Berechnung der nichtrationalen Funktion $y = \sqrt[3]{(1+x)(x^2-4)}$ das Radizieren, eine Rechenvorschrift, die im allgemeinen in endlich vielen Schritten nicht zum Ziele führt.

Wir wollen nun im folgenden untersuchen, wie weit sich jede Funktion $y = f(x)$, insbesondere jede transzendente Funktion, durch rationale Prozesse exakt ausdrücken oder wenigstens innerhalb einer vorgeschriebenen Genauigkeit annähern läßt. Ist dies möglich, so sagt man, die Funktion $y = f(x)$ läßt sich **approximieren**.

Bei ganzen rationalen Funktionen ist die Berechnung besonders einfach, deshalb geht man bei den Untersuchungen, die wir in den Abschnitten 2a und 2b durchführen werden, von den ganzen rationalen Funktionen aus.

1. Die Mittelwertsätze der Differentialrechnung

Zunächst wollen wir im Abschnitt 1 einige Sätze herleiten, die wir bei den Betrachtungen benötigen.

Satz 1:

Ist die Funktion $\varphi(x)$ im Intervall $a \leq x \leq b$ stetig und differenzierbar, ist ferner $\varphi(a) = \varphi(b)$, so gibt es mindestens eine Stelle ξ im Innern des Intervalls ($a < \xi < b$), für die $\varphi'(\xi) = 0$ ist.

Beweis:

Ist $\varphi(x)$ eine konstante Funktion, also $\varphi(x) = c$, so gilt für alle Werte x des Intervalls $\varphi'(x) = 0$. Also ist die Behauptung in diesem Falle richtig.

Ist dagegen die Funktion $\varphi(x)$ nicht konstant (vgl. Abb. 6), so gibt es im Innern des Intervalls mindestens eine Stelle ξ , an der die Funktion einen Extremwert annimmt¹⁾. Aus der Differenzierbarkeit der Funktion $\varphi(x)$ im Intervall folgt, daß an der Stelle ξ die Ableitung $\varphi'(\xi)$ existiert. An der Stelle eines Extremwertes muß aber die Ableitung notwendig gleich Null sein. Also gilt

$$\varphi'(\xi) = 0.$$

Damit ist der Satz, der als **Satz von Rolle** bezeichnet wird, bewiesen.

Satz 2:

Ist $f(x)$ im Intervall $a \leq x \leq b$ stetig und differenzierbar, so gibt es im Innern des Intervalls mindestens eine Stelle ξ , für die $\frac{f(b)-f(a)}{b-a} = f'(\xi)$ ist. (1)

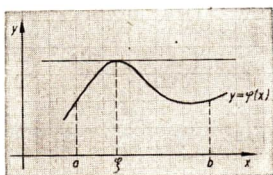


Abb. 6

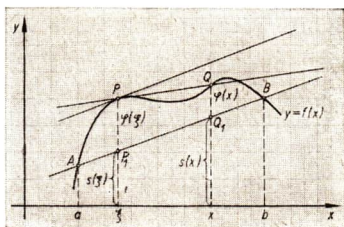


Abb. 7

Beweis:

Für $x = a$ und $x = b$ nimmt die Funktion $f(x)$ die Werte $f(a)$ bzw. $f(b)$ an (vgl. Abb. 7). Wir bilden nun eine Hilfsfunktion $\varphi(x)$, die die Differenz zwischen den Ordinaten der Funktion $f(x)$ und den Ordinaten der Sekante durch die Punkte $A(a; f(a))$ und $B(b; f(b))$ für die gleiche Abszisse x angibt. Bezeichnet man die Ordinate eines beliebigen Punktes auf der Sekante mit $s(x)$, so hat nach der Zweipunktegleichung der Geraden die Gleichung der Sekante die Form

$$\frac{f(b) - f(a)}{b - a} = \frac{s(x) - f(a)}{x - a}.$$

Daraus folgt die Beziehung

$$s(x) = f(a) + \frac{f(b) - f(a)}{b - a} (x - a).$$

¹⁾ Dies folgt aus der Stetigkeit der Funktion im abgeschlossenen Intervall

Bilden wir die Differenz, so erhalten wir

$$\varphi(x) = f(x) - s(x) = f(x) - f(a) - \frac{f(b) - f(a)}{b - a} (x - a).$$

Diese Funktion $\varphi(x)$ erfüllt hinsichtlich Stetigkeit und Differenzierbarkeit dieselben Bedingungen wie die Funktion $f(x)$. Außerdem ist

$$\varphi(a) = \varphi(b).$$

Also gibt es nach dem Satz von Rolle ein ξ mit $a < \xi < b$ derart, daß $\varphi'(\xi) = 0$ ist. Differenzieren wir $\varphi(x)$, so erhalten wir

$$\varphi'(x) = f'(x) - \frac{f(b) - f(a)}{b - a}.$$

Wegen $\varphi'(\xi) = 0$ folgt daraus

$$\frac{f(b) - f(a)}{b - a} = f'(\xi).$$

Damit ist die Behauptung bewiesen.

Dieser Satz wird als **Mittelwertsatz der Differentialrechnung** bezeichnet.

Der Beweis kann noch auf eine andere Weise geführt werden: Wir bilden die Hilfsfunktion

$$g(x) = f(x) + \lambda x,$$

wobei λ eine noch zu bestimmende Konstante ist. Da $f(x)$ differenzierbar ist, besitzt auch $g(x)$ diese Eigenschaft. Wir wollen nun λ so bestimmen, daß $g(a) = g(b)$ ist. Dann gilt

$$f(a) + \lambda a = f(b) + \lambda b.$$

Daraus folgt

$$\lambda = -\frac{f(b) - f(a)}{b - a}.$$

Mit dieser Festsetzung für λ ist $g(x)$ eine Funktion, die alle Voraussetzungen des Satzes von Rolle erfüllt. Es gibt also ein ξ derart, daß $g'(\xi) = 0$ ist. Die Ableitung von $g(x)$ ergibt sich zu

$$g'(x) = f'(x) + \lambda.$$

Für $x = \xi$ gilt also

$$0 = f'(\xi) + \lambda.$$

Setzen wir in diese letzte Beziehung den vorher bestimmten Wert von λ ein, so erhalten wir

$$0 = f'(\xi) - \frac{f(b) - f(a)}{b - a}$$

oder

$$\frac{f(b) - f(a)}{b - a} = f'(\xi).$$

Der linke Ausdruck dieser Relation, der Differenzenquotient, ist geometrisch der Anstieg der Sekante, die durch die Punkte $A[a; f(a)]$ und $B[b; f(b)]$ geht. Der Ausdruck $f'(\xi)$ ist der Anstieg der Tangente im Berührungspunkt $P[\xi; f(\xi)]$.

Dieses Beweisverfahren mit Hilfe der Konstanten λ werden wir beim Satz 3 anwenden.

Aus dem Mittelwertsatz folgt

$$f(b) = f(a) + (b - a) f'(\xi) \quad \text{mit} \quad a < \xi < b. \quad (2)$$

Ersetzt man in (2) $b - a$ durch h und a durch x_0 , so ergibt sich

$$f(x_0 + h) = f(x_0) + h \cdot f'(\xi).$$

Wegen $x_0 < \xi < x_0 + h$ kann man für ξ auch $x_0 + \vartheta h$ schreiben, wobei ϑ eine Zahl zwischen 0 und 1 ist, die man bei der Anwendung des Mittelwertsatzes oft nicht näher bestimmen kann, die man aber — wie sich herausstellen wird — auch nicht näher zu kennen braucht. Für $f(x_0 + h)$ ergibt sich dann

$$f(x_0 + h) = f(x_0) + h f'(x_0 + \vartheta h) \quad \text{mit} \quad 0 < \vartheta < 1. \quad (3)$$

Diese Form des Mittelwertsatzes gibt eine erste Annäherung an die Funktion $f(x)$ in der Umgebung der Stelle $x = x_0$.

Wenn die Ableitung $f'(x)$ eine monotone Funktion ist, so kann man mit (3) die Funktionswerte $f(x_0 + h)$ abschätzen, wenn man den Funktionswert $f(x_0)$ kennt. In diesem Falle liegt nämlich der größte bzw. der kleinste Wert von $f'(x)$ in dem einen oder anderen Endpunkt des Intervalls. Man kann also für ϑ die Werte 0 oder 1 als untere oder obere Schranke einsetzen.

1. Beispiel:

$\ln 10 = 2,3026$; wie groß ist $\ln 10,2$?

$$x = 10, \quad h = 0,2; \quad \ln 10,2 = \ln x + \frac{h}{x + \vartheta h}.$$

$$\text{Wir setzen } \vartheta = 1, \quad \frac{h}{x + \vartheta h} = \frac{0,2}{10,2} = 0,0196$$

$$\text{und } \vartheta = 0, \quad \frac{h}{x + \vartheta h} = \frac{0,2}{10} = 0,02.$$

Folglich gilt die Abschätzung $2,3226 > \ln 10,2 > 2,3222$.

2. Beispiel:

$e = 2,718$; wie groß ist $e^{1,1}$?

$$x = 1, \quad h = 0,1; \quad e^{x+h} = e^x + h \cdot e^{x+\vartheta h}.$$

$$\text{Wir setzen } \vartheta = 0: e^{1,1} > e^1 + 0,1 \cdot e^1 = 1,1 \cdot e^1 = 2,9898;$$

$$\vartheta = 1: e^{1,1} < e^1 + 0,1 \cdot e^{1,1}; \quad 0,9 \cdot e^{1,1} < e^1; \quad e^{1,1} < \frac{e}{0,9} = 3,0200;$$

also gilt die Abschätzung $2,9898 < e^{1,1} < 3,0200$.

Satz 3:

Wenn $f(x)$ und $g(x)$ im Intervall $a \leq x \leq b$ differenzierbare Funktionen sind und $g'(x)$ an keiner Stelle des Intervalls $a < x < b$ gleich Null ist, gilt

$$\frac{f(x+h)-f(x)}{g(x+h)-g(x)} = \frac{f'(x+\vartheta h)}{g'(x+\vartheta h)} \quad \text{mit } 0 < \vartheta < 1.$$

Beweis:

Wir bilden eine Hilfsfunktion $\varphi(x)$, die der Bedingung

$$\varphi(x) = f(x) + \lambda g(x)$$

genügt. Dabei sei λ ein Wert, den wir so bestimmen, daß $\varphi(a) = \varphi(b)$ ist. Aus den Gleichungen

$$\varphi(a) = f(a) + \lambda g(a) \quad \text{und} \quad \varphi(b) = f(b) + \lambda g(b)$$

folgt wegen der Bedingung $\varphi(a) = \varphi(b)$ die Relation

$$f(a) + \lambda g(a) = f(b) + \lambda g(b).$$

Daraus ergibt sich für λ der Wert

$$\lambda = -\frac{f(b) - f(a)}{g(b) - g(a)}.$$

Wenden wir dagegen auf die Funktion $\varphi(x)$ den Satz von Rolle an, so erhalten wir wegen

$$\varphi'(x) = f'(x) + \lambda g'(x)$$

die Beziehung

$$\varphi'(\xi) = f'(\xi) + \lambda g'(\xi) = 0.$$

Damit erhalten wir auf eine andere Weise für λ den Wert

$$\lambda = -\frac{f'(\xi)}{g'(\xi)}.$$

Das heißt, es gilt die Gleichung

$$\frac{f(b) - f(a)}{g(b) - g(a)} = \frac{f'(\xi)}{g'(\xi)} \quad \text{mit } a < \xi < b.$$

Es ist üblich, in dieser Gleichung a durch x und b durch $x+h$ mit $h > 0$ und folglich ξ durch $x+\vartheta h$ mit $0 < \vartheta < 1$ zu ersetzen. Man erhält dann

$$\frac{f(x+h) - f(x)}{g(x+h) - g(x)} = \frac{f'(x+\vartheta h)}{g'(x+\vartheta h)}.$$

Dabei ist h ein Zuwachs der Variablen x . Zum gleichen Ergebnis gelangt man, wenn man a durch $x+h$ mit $h < 0$ und b durch x ersetzt.

Damit ist die Gültigkeit des Satzes 3 für jedes h erwiesen. Man bezeichnet den Satz als **verallgemeinerten Mittelwertsatz**.

2. Die Taylorsche Sätze

a) Die Entwicklung einer ganzen rationalen Funktion an der Stelle $x = 0$

Wie wir wissen (vgl. Lehrbuch der Mathematik, 11. Schuljahr, Seite 7), läßt sich jede ganze rationale Funktion n -ten Grades in der Form

$$y = G(x) = a_0 + a_1 x + a_2 x^2 + a_3 x^3 + \cdots + a_n x^n \quad (4)$$

schreiben. Dabei sind die Koeffizienten a_k ($k = 0, 1, 2, \dots, n-1$) beliebige reelle Zahlen, und es ist a_n eine von Null verschiedene reelle Zahl.

Wir wollen eine Beziehung zwischen den Koeffizienten a_k der Funktion einerseits und den Werten der Funktion sowie ihrer Ableitungen an der Stelle $x = 0$ andererseits herstellen. Zu diesem Zweck bilden wir die ersten n Ableitungen der Funktion und erhalten

$$G'(x) = a_1 + 2a_2 x + 3a_3 x^2 + \cdots + n \cdot a_n x^{n-1}$$

$$G''(x) = 2a_2 + 3 \cdot 2 \cdot a_3 x + \cdots + n \cdot (n-1) \cdot a_n x^{n-2}$$

$$G'''(x) = 3 \cdot 2 \cdot a_3 + \cdots + n \cdot (n-1) \cdot (n-2) \cdot a_n x^{n-3}$$

$$G^{(n)}(x) = n \cdot (n-1) \cdot (n-2) \cdot \cdots \cdot 3 \cdot 2 \cdot 1 \cdot a_n = n! \cdot a_n.$$

Setzen wir in diesen $(n+1)$ Gleichungen $x = 0$ und verwenden wir für das Produkt $1 \cdot 2 \cdot 3 \cdot \cdots \cdot k$ die Abkürzung $k!$ (mit $0! = 1$), so erhalten wir für die Koeffizienten a_k ($k = 0, 1, 2, \dots, n$) die Beziehungen

$$G^{(k)}(0) = k! a_k \quad \text{oder} \quad a_k = \frac{1}{k!} G^{(k)}(0), \quad (k = 0, 1, 2, \dots, n).$$

Für $k = 0$ setzt man sinngemäß $G^{(0)}(x) = G(x)$, das heißt, die nullte Ableitung einer Funktion wird als die Funktion selbst definiert. Somit läßt sich die ganze rationale Funktion $G(x)$ darstellen in der Form

$$y = G(x) = G(0) + \frac{x}{1!} G'(0) + \frac{x^2}{2!} G''(0) + \frac{x^3}{3!} G'''(0) + \cdots + \frac{x^n}{n!} G^{(n)}(0). \quad (5)$$

Wir sagen: Die Formel (5) stellt die Entwicklung der Funktion $G(x)$ an der Stelle $x_0 = 0$ dar. Sie gibt also an, wie eine beliebige ganze rationale Funktion nach steigenden Potenzen von x geordnet werden kann.

Beispiel:

$$y = G(x) = (x+1)^6 - (2-x)^4 + 7(x+2).$$

Entwicklungsgang:

k	$G^{(k)}(x)$	$G^{(k)}(0)$	$k!$	a_k
0	$(x+1)^6 - (2-x)^4 + 7(x+2)$	-1	1	-1
1	$6(x+1)^5 + 4(2-x)^3 + 7$	+45	1	+45
2	$30(x+1)^4 - 12(2-x)^2$	-18	2	-9
3	$120(x+1)^3 + 24(2-x)$	+168	6	+28
4	$360(x+1)^2 - 24$	+336	24	+14
5	$720(x+1)$	+720	120	+6
6	720	+720	720	+1

Es ist also

$$y = -1 + 45x - 9x^2 + 28x^3 + 14x^4 + 6x^5 + x^6.$$

b) Die Entwicklung einer ganzen rationalen Funktion an der Stelle $x = a$

Wie uns bekannt ist (vgl. Lehrbuch der Mathematik, 9. Schuljahr, Seite 24), läßt sich die Funktion zweiten Grades $y = ax^2 + bx + c$ umformen zu $y = a \cdot (x - d)^2 + e$. Das bedeutet nichts anderes, als daß die Entwicklung nach Potenzen von x in eine solche nach Potenzen von $(x - d)$ überführt worden ist. Eine Umformung dieser Art ist für jede ganze rationale Funktion an jeder beliebigen Stelle $x = a$ möglich. Wir führen für die Funktion

$$G(x) = c_0 + c_1(x - a) + c_2(x - a)^2 + c_3(x - a)^3 + \cdots + c_n(x - a)^n \quad (6)$$

die Substitution $x - a = \xi$ durch. Das bedeutet geometrisch die Parallelverschiebung der Ordinate in den Abszissenpunkt $x - a$. Eine Entwicklung an der Stelle $x = a$ ist also gleichbedeutend mit einer Entwicklung an der Stelle $\xi = 0$. Damit ist das Problem auf den in dem Abschnitt 2a behandelten Fall zurückgeführt.

Wir wollen nun die noch unbekannten Koeffizienten c_k der Entwicklung (6) bestimmen. Bilden wir wieder die Ableitung $G^{(k)}(x)$ und setzen sodann $x = a$, so erhalten wir die Beziehungen

$$G^{(k)}(a) = k!c_k \quad \text{oder} \quad c_k = \frac{1}{k!} G^{(k)}(a), \quad (k = 0, 1, 2, \dots, n).$$

Damit folgt aus dem Ansatz (6)

$$G(x) = G(a) + \frac{(x-a)}{1!} G'(a) + \frac{(x-a)^2}{2!} G''(a) + \cdots + \frac{(x-a)^n}{n!} G^{(n)}(a). \quad (7)$$

Die Darstellung der Funktion $G(x)$ in der Form (7) ist mit der Darstellung

$$G(x) = c_0 + c_1(x - a) + \cdots + c_n(x - a)^n$$

identisch.

Ersetzt man a durch x_0 , so erhält man eine häufig verwendete Darstellung:

$$G(x) = G(x_0) + \frac{(x-x_0)}{1!} G'(x_0) + \frac{(x-x_0)^2}{2!} G''(x_0) + \cdots + \frac{(x-x_0)^n}{n!} G^{(n)}(x_0). \quad (8)$$

Wir wollen bei unseren folgenden Überlegungen stets x_0 an Stelle von a verwenden.

Wir sagen in diesem Falle, die Funktion läßt sich an der Stelle $x_0 = a$ entwickeln. Für $x_0 = 0$ geht die Formel (8) in die Formel (5) über. Ersetzt man x durch $x_0 + h$, wobei x_0 konstant und h variabel ist, so erhält man schließlich

$$G(x_0 + h) = G(x_0) + \frac{h}{1!} G'(x_0) + \frac{h^2}{2!} G''(x_0) + \cdots + \frac{h^n}{n!} G^{(n)}(x_0). \quad (9)$$

Diese Entwicklung der ganzen rationalen Funktion $G(x)$ an der Stelle x_0 heißt die **Taylorische Formel** für $G(x)$. Somit sind auch die Darstellungen (8) und (5) Taylorische Formeln. Die Entwicklung (5) ist ein Spezialfall der Taylorischen Formel mit $x_0 = 0$ und $h = x$ und wird auch als **Maclaurinsche Formel** bezeichnet.

Jede ganze rationale Funktion einer Veränderlichen, das heißt jedes Polynom in einer Veränderlichen, kann als Taylorische Formel geschrieben werden.

Beispiel:

Die Funktion $G(x) = \frac{1}{3} \cdot x(x^2 - 3)$ ist an den Stellen $x_0 = 0, 1, 2$ zu entwickeln!

k	$G^{(k)}(x)$	$k!$	$G^{(k)}(0)$	a_k	$G^{(k)}(1)$	c_k	$G^{(k)}(2)$	c_k
0	$\frac{1}{3} \cdot x(x^2 - 3)$	1	0	0	$-\frac{2}{3}$	$-\frac{2}{3}$	$+\frac{2}{3}$	$+\frac{2}{3}$
1	$x^2 - 1$	1	-1	-1	0	0	+3	+3
2	$2x$	2	0	0	+2	+1	+4	+2
3	2	6	+2	$+\frac{1}{3}$	+2	$+\frac{1}{3}$	+2	$+\frac{1}{3}$

Somit erhält man die folgenden Entwicklungen:

$$\text{Für } x_0 = 0: \quad y = \quad - \quad x \quad + \frac{1}{3} x^3,$$

$$\text{für } x_0 = 1: \quad y = -\frac{2}{3} + \quad (x-1)^2 + \frac{1}{3} (x-1)^3,$$

$$\text{für } x_0 = 2: \quad y = +\frac{2}{3} + 3(x-2) + 2(x-2)^2 + \frac{1}{3} (x-2)^3,$$

für die man auch durch Ausrechnen nachweisen kann, daß alle drei Entwicklungen die gleiche Funktion darstellen.

c) Approximation und Schmiegunparabeln

Die Taylorsche Formel ermöglicht die näherungsweise Bestimmung der Funktionswerte an der Stelle $x = x_0$. Ist die Differenz $x - x_0$ genügend klein, so geben die höheren Potenzen dieser Differenzen einen verhältnismäßig kleinen Beitrag zum Funktionswert. Vernachlässigt man diese Glieder, so erhält man einen Wert, der zwar vom exakten Funktionswert abweicht, aber dennoch eine hinreichende Näherung für diesen Funktionswert darstellt. Man nennt diese Entwicklung **Approximation** der gegebenen Funktion. Alle Entwicklungen, die die gegebene Funktion näherungsweise darstellen, heißen **approximierende Funktionen** der gegebenen Funktion.

Wir wollen die Funktion

$$y = \frac{1}{3} x(x^2 - 3)$$

an der Stelle $x_0 = 2$ entwickeln:

$$y = \frac{2}{3} + 3(x - 2) + 2(x - 2)^2 + \frac{1}{3}(x - 2)^3.$$

Für die Stelle $x_1 = 2,1$ soll y auf zwei Dezimalen genau berechnet werden. Man erkennt, daß das vierte Glied $\frac{1}{3}(x - 2)^3$ für $x_1 = 2,1$ kleiner als 0,001 ist, das heißt unter den gegebenen Bedingungen genügt die Darstellung bis zur zweiten Potenz. Ist $x_2 = 2,02$ und wird für y wieder eine Genauigkeit auf zwei Dezimalen gefordert, so reichen zur Bestimmung des Funktionswertes schon die ersten beiden Glieder der Darstellung aus, da das quadratische und das kubische Glied kleiner als 0,001 sind.

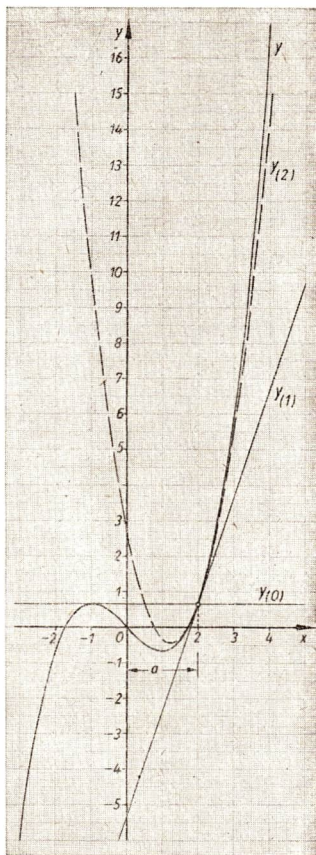


Abb. 8

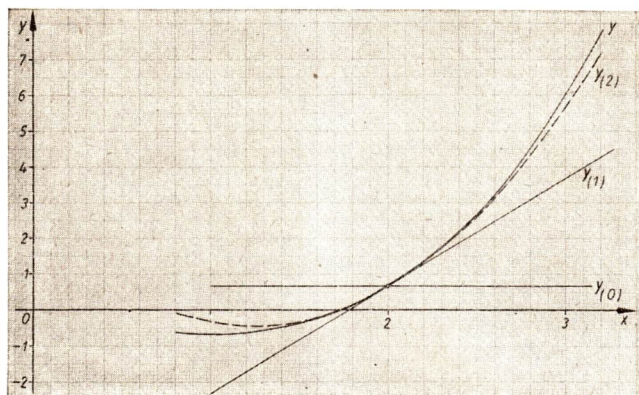


Abb. 9

In unserem Beispiel ergeben sich die approximierenden Funktionen zu

$$\begin{aligned}
 y_{(0)} &= \frac{2}{3} & &= \frac{2}{3}, \\
 y_{(1)} &= \frac{2}{3} + 3(x-2) & &= 3x - \frac{16}{3}, \\
 y_{(2)} &= \frac{2}{3} + 3(x-2) + 2(x-2)^2 = 2x^2 - 5x + \frac{8}{3}.
 \end{aligned}$$

Man nennt den approximierenden Ausdruck $y_{(k)}$ vom Grade k eine Approximation k -ter Ordnung zur Funktion $y = G(x)$.

Den approximierenden Funktionen entsprechen bestimmte Kurven. In unserem Beispiel ist:

$y_{(0)}$ eine Parallele zur x -Achse durch den Punkt $\left(2; \frac{2}{3}\right)$,

$y_{(1)}$ eine Gerade durch den Punkt $\left(2; \frac{2}{3}\right)$ mit dem Anstieg $m = 3$,

$y_{(2)}$ eine Parabel, deren Achse zur y -Achse parallel ist und die den Scheitel $\left(+\frac{5}{4}; -\frac{11}{24}\right)$ hat. Sie geht durch den Punkt $\left(2; \frac{2}{3}\right)$.

Die Funktion selbst ist eine kubische Parabel mit den Nullstellen $x_1 = 0$; $x_{2,3} = \pm\sqrt[3]{3}$ und den Extremstellen $x_E = \pm 1$.

Die Kurven der approximierenden Funktionen nullter bis zweiter Ordnung und die Funktion selbst sind in der Abbildung 8 dargestellt. Deutlich ist zu erkennen, daß die Approximation um so besser ist, je näher man sich bei der Stelle $x_0 = 2$ befindet und je höher die Ordnung der Approximation ist.

Man nennt die Kurven der approximierenden Funktionen die **Schmiegungsparabeln** an der Stelle $x_0 = a$ für die durch die gegebene Funktion erklärte Kurve. Der Schmiegungscharakter tritt noch deutlicher in der Abbildung 9 hervor, die einen Ausschnitt aus der Abbildung 8 gibt, in dem der Maßstab der x -Achse im Verhältnis 5 : 1 vergrößert ist.

Wir haben erkannt, daß eine ganze rationale Funktion durch die Taylorsche Formel in beliebig vielen, durch die Wahl von x_0 bedingten Entwicklungen darstellbar ist. Die Koeffizienten der Darstellung sind durch die Werte der Funktion und ihrer Ableitungen an der Stelle $x_0 = 0$ bzw. $x_0 = a$ bestimmt.

d) Der Taylorsche Lehrsatz für eine beliebige Funktion

Wie wir bereits erkannt haben (vgl. Abschnitt 2, b), ist für jede ganze rationale Funktion vom Grade $(n - 1)$, also für

$$G(x) = a_0 + a_1 x + \cdots + a_{n-1} x^{n-1}$$

die Entwicklung

$$G(x_0 + h) = G(x_0) + \frac{h}{1!} G'(x_0) + \frac{h^2}{2!} G''(x_0) + \cdots + \frac{h^{(n-1)}}{(n-1)!} G^{(n-1)}(x_0)$$

möglich.

Wir wollen jetzt untersuchen, ob und inwieweit eine ähnliche Entwicklung für Funktionen möglich ist, die keine ganzen rationalen sind.

Es sei $f(x)$ irgendeine Funktion, die an der Stelle x_0 und in einer gewissen Umgebung derselben mindestens n -mal differenzierbar ist. Für $n \geq 1$ ist dann die Funktion $f(x)$ gewiß an der Stelle x_0 und in einer gewissen Umgebung derselben definiert. Es sei $h \neq 0$ seinem Betrage nach so klein, daß $x_0 + h$ dieser Umgebung angehört. Wir wollen nun untersuchen, inwieweit der Funktionswert $f(x_0 + h)$ durch den Ausdruck

$$f(x_0) + \frac{f'(x_0)}{1!} h + \frac{f''(x_0)}{2!} h^2 + \cdots + \frac{f^{(n-1)}(x_0)}{(n-1)!} h^{n-1} \quad (9a)$$

dargestellt wird. Die folgenden Überlegungen werden ergeben, daß diese beiden Funktionswerte für beliebige Funktionen im allgemeinen nicht gleich sind, daß aber der Unterschied unter gewissen, sehr allgemeinen Voraussetzungen seinem Betrage nach so klein ist, daß der Ausdruck (9a) als ein angenäherter Wert für $f(x_0 + h)$ angesehen werden kann.

Um für $f(x_0 + h)$ einen genauen Wert mit Hilfe der Entwicklung (9a) zu erhalten, müssen wir die Differenz zwischen beiden Funktionswerten ausgleichen, das heißt wir müssen sie dem Ausdruck (9a) zufügen. Diese Differenz wird Restglied genannt und mit R_n bezeichnet. Der Index n soll dabei zum Ausdruck bringen, daß das Restglied nach dem n -ten Summanden hinzugefügt wird.

Wir behaupten nun den

Satz 4:

Ist $f(x)$ eine im Intervall $a \leq x \leq b$ stetige und n -mal differenzierbare Funktion, so ist für die Funktion die Entwicklung

$$f(x_0 + h) = f(x_0) + \frac{h}{1!} f'(x_0) + \frac{h^2}{2!} f''(x_0) + \dots + \frac{h^{(n-1)}}{(n-1)!} f^{(n-1)}(x_0) + R_n \quad (10)$$

mit

$$R_n = \frac{h^n}{n!} f^{(n)}(x_0 + \vartheta h) \quad 0 < \vartheta < 1$$

möglich.

Der wesentliche Inhalt dieses Satzes besteht nicht darin, daß ein Ansatz der Form (10) überhaupt gemacht werden kann. Das Schwergewicht dieses Problems liegt ausschließlich in den Aussagen, die über das Restglied gemacht werden können.

Beweis des Satzes 4:

Wir müssen also die in Satz 4 über R_n gemachten Aussagen beweisen. Wir werden dazu den erweiterten Mittelwertsatz auf zwei passend gewählte Hilfsfunktionen anwenden.

Wir gehen aus von der Beziehung

$$R_n = f(x_0 + h) - f(x_0) - \frac{f'(x_0)}{1!} h - \frac{f''(x_0)}{2!} h^2 - \dots - \frac{f^{(n-1)}(x_0)}{(n-1)!} h^{n-1}. \quad (11)$$

In dieser Relation ersetzen wir $x_0 + h$ durch b , x_0 durch a und folglich h durch $b - a$. Dann erhalten wir

$$R_n = f(b) - f(a) - \frac{b-a}{1!} f'(a) - \frac{(b-a)^2}{2!} f''(a) - \dots - \frac{(b-a)^{n-1}}{(n-1)!} f^{(n-1)}(a). \quad (11a)$$

Wir setzen nun an Stelle der Konstanten a die Veränderliche x und erhalten damit die erste der zu betrachtenden Hilfsfunktionen, die wir mit $\varphi(x)$ bezeichnen wollen:

$$\varphi(x) = f(b) - f(x) - \frac{b-x}{1!} f'(x) - \frac{(b-x)^2}{2!} f''(x) - \dots - \frac{(b-x)^{n-1}}{(n-1)!} f^{(n-1)}(x). \quad (12)$$

Diese Funktion hat folgende Eigenschaften:

- Sie ist in dem Intervall $a \leq x \leq b$ mindestens einmal differenzierbar, da sie sich aus Gliedern zusammensetzt, die ein Produkt aus einer ganzen rationalen Funktion in x und einer Ableitung von $f(x)$ von höchstens $(n-1)$ -ter Ordnung sind. (Nach Voraussetzung ist aber $f(x)$ im Intervall $a \leq x \leq b$ n -mal differenzierbar.)
- Ihre Ableitung $\varphi'(x)$ ergibt sich zu

$$\begin{aligned} \varphi'(x) = & -f'(x) + f'(x) - \frac{b-x}{1!} f''(x) + \frac{b-x}{1!} f''(x) - \frac{(b-x)^2}{2!} f'''(x) \dots \\ & + \frac{(b-x)^{n-2}}{(n-2)!} f^{(n-1)}(x) - \frac{(b-x)^{n-1}}{(n-1)!} f^{(n)}(x). \end{aligned}$$

¹⁾ Wir wollen uns an zwei Gliedern von $\varphi(x)$ klarmachen, wie der einfache Ausdruck für $\varphi'(x)$ entsteht. Nach der Produktregel ergibt sich für $-\frac{(b-x)^3}{3!} f^{(3)}(x) - \frac{(b-x)^4}{4!} f^{(4)}(x)$ die Ableitung zu $+\frac{(b-x)^2}{2!} f^{(3)}(x) - \frac{(b-x)^3}{3!} f^{(4)}(x) + \frac{(b-x)^3}{3!} f^{(4)}(x) - \frac{(b-x)^4}{4!} f^{(5)}(x) -$
 $+\frac{(b-x)^2}{2!} f^{(3)}(x) - \frac{(b-x)^4}{4!} f^{(4)}(x).$

Wir erkennen, daß sich alle Glieder bis auf das letzte wegheben, so daß wir schließlich

$$\varphi'(x) = - \frac{(b-x)^{n-1}}{(n-1)!} f^{(n)}(x)$$

erhalten.

c) Es ist $\varphi(a) = R_n$ und $\varphi(b) = 0$ ¹⁾.

Als zweite Hilfsfunktion benötigen wir eine Funktion mit folgenden Eigenschaften:

- a) Die Funktion muß ebenfalls im Intervall $a \leq x \leq b$ differenzierbar sein. Ihre Ableitung muß für $a < x < b$ stets von Null verschieden sein.
- b) An der Stelle b soll die Funktion gleich Null, an der Stelle a soll sie verschieden von Null sein.

Die Funktion

$$\psi(x) = (b-x)^n$$

erfüllt diese Bedingungen. Für sie ist $\psi(a) = (b-a)^n$ und $\psi'(x) = -n(b-x)^{n-1}$. Dabei soll n die gleiche natürliche Zahl wie in $\varphi(x)$ sein.

Wenden wir auf beide Hilfsfunktionen den verallgemeinerten Mittelwertsatz an, so erhalten wir

$$\frac{\varphi(b) - \varphi(a)}{\psi(b) - \psi(a)} = \frac{\varphi'(\xi)}{\psi'(\xi)} \quad \text{mit} \quad a < \xi < b. \quad (13)$$

Es gilt also

$$\frac{0 - R_n}{0 - (b-a)^n} = \frac{\frac{(b-\xi)^{n-1}}{(n-1)!} f^{(n)}(\xi)}{n(b-\xi)^{n-1}} \quad (13a)$$

oder

$$R_n = \frac{(b-a)^n f^{(n)}(\xi)}{n!}. \quad (14)$$

Wegen $a < \xi < b$ können wir

$$\xi = a + \vartheta(b-a)$$

setzen, wobei $0 < \vartheta < 1$ ist. Dann erhalten wir

$$R_n = \frac{(b-a)^n}{n!} f^{(n)}[a + \vartheta(b-a)]. \quad (14a)$$

Setzen wir nun für a und b wieder die ursprünglichen Bezeichnungen x_0 und $x_0 + h$ ein, so ergibt sich schließlich

$$R_n = \frac{h^n}{n!} f^{(n)}(x_0 + \vartheta h). \quad (14b)$$

¹⁾ Diese Werte ergeben sich, wenn man in $\varphi(x)$ für $x = a$ bzw. $x = b$ setzt.

Es gilt also

$$f(x_0 + h) = f(x_0) + \frac{h}{1!} f'(x_0) + \frac{h^2}{2!} f''(x_0) + \dots + \frac{h^{(n-1)}}{(n-1)!} f^{(n-1)}(x_0) + \frac{h^n}{n!} f^{(n)}(x_0 + \vartheta h). \quad (10a)$$

Die Entwicklung (10a) der Funktion $f(x)$ heißt der **Taylorische Lehrsatz** für die Funktion $f(x)$. Sie enthält die Taylorische Formel für ganze rationale Funktionen. Den Ausdruck $R_n = \frac{h^n}{n!} f^{(n)}(x_0 + \vartheta h)$ bezeichnet man als **Restglied von Lagrange**.

Die Entwicklung

$$f(x_0 + h) = f(x_0) + \frac{h}{1!} f'(x_0) + \dots + \frac{h^{n-1}}{(n-1)!} f^{(n-1)}(x_0) + R_n$$

ist immer dann möglich, wenn alle Ableitungen bis zur $(n-1)$ -ten an der Stelle x_0 existieren. Man erhält für die n ersten Summanden der rechten Seite wohlbestimmte Werte, jedoch kann man über das Restglied R_n keine Aussage machen. Aus der Formel erkennen wir, daß die n ersten Summanden der rechten Seite einen Näherungswert für den Funktionswert $f(x_0 + h)$ geben (n -te approximierende Funktion) und daß das Restglied R_n die Differenz zwischen dem Funktionswert $f(x_0 + h)$ und diesem Näherungswert ist. Die Bedeutung des Taylorschen Lehrsatzes liegt darin, daß die Größe des Restgliedes

$$R_n = \frac{h^n}{n!} f^{(n)}(x_0 + \vartheta h) \quad (0 < \vartheta < 1)$$

einer so entwickelten Funktion $f(x)$ unter gewissen Bedingungen abgeschätzt werden kann. Zum Beispiel ist die Abschätzung immer dann möglich, wenn im Intervall $a \leq x \leq b$ die n ersten Ableitungen existieren und die n -te Ableitung der Funktion monoton ist. Aus der Größenordnung des Restgliedes kann man dann schließen, ob die erzielte Annäherung für den jeweiligen Zweck ausreicht.

Aufgaben

1. Entwickeln Sie die folgenden Funktionen an der Stelle $x = 0$!

a) $y = (x^2 - 1)^3 + (x^3 - 2x - 4)^2$

b) $y = (x^2 - 2)^3 - (x^3 + 2x - 4)^3$

2. Für eine ganze rationale Funktion n -ten Grades $y = G(x)$ sei

a) $n = 3$; $G(0) = 1$, $G'(0) = 2$, $G''(0) = \frac{1}{2}$, $G'''(0) = 2$;

b) $n = 4$; $G(0) = 2$, $G'(0) = 0$, $G''(0) = 2$, $G'''(0) = 0$, $G^{(4)}(0) = 3$;

c) $n = 5$; $G(0) = G'(0) = G''(0) = G'''(0) = 0$, $G^{(4)}(0) = 8$, $G^{(5)}(0) = 3$.

Es ist die Polynomdarstellung der Funktion anzugeben.

3. Lesen Sie aus der Polynomdarstellung

a) $y = 1 - \frac{1}{2}x - \frac{1}{2}x^2 - x^3$, b) $y = 2x - 4x^2 - \frac{1}{12}x^3$,

c) $y = x^3 - \frac{1}{12}x^5$, d) $y = 1 - x^2$

die Werte der Funktion und ihrer Ableitungen an der Stelle $x = 0$ ab!

4. Die folgenden Funktionen sind nach der Taylorsche Formel zu entwickeln:

a) $y = (1 + x)^2$,

b) $y = (1 + x)^4$,

c) $y = (1 + x)^3$,

d) $y = (1 + 2x)^4$,

e) $y = (1 + x^2)^3$,

f) $y = \left(1 + \frac{1}{2}x\right)^5$,

g) $y = (2 + x)^5$,

h) $y = (2 + 3x)^4$,

i) $y = \left(1 + \frac{1}{2}x^2\right)^4$

5. Geben Sie für die Funktion $y = \frac{1}{3}x(x^2 - 3)$ die Entwicklung an der Stelle $x = 1$ an!

Wie lauten die approximierenden Funktionen 0., 1. und 2. Ordnung? Stellen Sie Wertetabellen her und zeichnen Sie die Schmiegeparabeln!

6. Es ist die Funktion $y = \frac{1}{4}(x^2 - 1)(x^2 - 7)$ an der Stelle

a) $x = 0$, b) $x = 1$, c) $x = 2$, d) $x = 3$ zu entwickeln.

7. Es sind die approximierenden Funktionen zu der in Aufgabe 6 angeführten Funktion an der Stelle $x = 0$ zu bilden. Die Schmiegeparabeln sind zu zeichnen.

8. Entwickeln Sie für die gleiche Funktion wie in Aufgabe 6 die approximierenden Funktionen an der Stelle $x = 2$! Zeichnen Sie die Schmiegeparabeln 3. Ordnung!

9. Es ist die Parabel 2. Grades $y = f(x)$ anzugeben, die für $x = 2$ den Wert $y = 2$ und im Punkte $(2; 2)$ die Steigung $m = 1$ hat und für die $f''(2) = 2$ ist!

10. Die folgenden ganzen rationalen Funktionen n -ten Grades sind an der Stelle $x = a$ zu entwickeln:

a) $n = 3$; $f(1) = -1$, $f'(1) = \frac{1}{2}$, $f''(1) = \frac{1}{2}$, $f'''(1) = 2$; $a = 1$,

b) $n = 5$; $f(3) = 0$, $f'(3) = -1$, $f''(3) = f'''(3) = f^{(4)}(3) = 0$, $f^{(5)}(3) = 60$; $a = 3$,

c) $n = 4$; $f(-2) = 1$, $f'(-2) = f''(-2) = 0$, $f'''(-2) = 3$, $f^{(4)}(-2) = 8$; $a = -2$.

Aus den Entwicklungen ist die für $a = 0$ abzuleiten. Es sind die Werte der Funktion und ihrer Ableitungen für $a = 0$ abzulesen.

11. Ein gleichmäßig beschleunigter Körper bewegt sich auf einer Geraden. Er befindet sich zur Zeit $t = 0$ an der Stelle $y(0) = 10$ cm. Seine Geschwindigkeit für $t = 0$ ist

$$v(0) = \left(\frac{dy}{dt} \right)_{t=0} = 2 \frac{\text{cm}}{\text{sec}},$$

seine Beschleunigung zur gleichen Zeit

$$b(0) = \left(\frac{d^2 y}{dt^2} \right)_{t=0} = 1 \frac{\text{cm}}{\text{sec}^2}.$$

Wie lautet das Weg-Zeit-Gesetz $y = f(t)$ der Bewegung? Fertigen Sie das Weg-Zeit-Diagramm an!

12. Lösen Sie die gleiche Aufgabe für $t = 0$; $y(0) = 12$ m, $v(0) = 8 \frac{\text{m}}{\text{sec}}$, $b(0) = 2 \frac{\text{m}}{\text{sec}^2}$. Für welchen Wert von t ergibt sich $y = 0$?

13. Zur Zeit $t = 4$ sec hat ein gleichförmig beschleunigter Körper die Höhe $h(4) = 80$ m erreicht. Seine Geschwindigkeit an dieser Stelle ist $\frac{dh}{dt} = v(4) = 40 \frac{\text{m}}{\text{sec}}$ und seine Beschleunigung $\frac{d^2 h}{dt^2} = b(4) = 10 \frac{\text{m}}{\text{sec}^2}$. Es ist die Bewegungsfunktion $h = f(t)$ am Zeitpunkt $t = 4$ zu entwickeln. Die Funktion ist nach Potenzen von t zu ordnen. Um welche Bewegung handelt es sich?

14. Ein Körper bewegt sich in der x, y -Ebene. Die Komponenten der Bewegung parallel zu den Koordinatenachsen sind gleichmäßig beschleunigte Bewegungen. Zur Zeit $t = 2$ ergeben sich die folgenden Beziehungen:

für die x -Komponente $x = g(t)$: $g(2) = 4$, $g'(2) = -2$, $g''(2) = -2$;

für die y -Komponente $y = h(t)$: $h(2) = 1$, $h'(2) = 4$, $h''(2) = +2$.

a) Die Funktionen $x = g(t)$, $y = h(t)$ sind zu bestimmen.

b) Wie lautet die Funktion $y = f(x)$ der Bahnkurve des bewegten Körpers?

Anleitung:

Es sind die Gleichungen für x und y zu addieren. Daraus ist die Größe $(t-2)$ zu berechnen. Das so gefundene Ergebnis ist in die erste Gleichung einzusetzen und die Gleichung nach y aufzulösen.

c) Es ist die Bahnkurve graphisch darzustellen und die Kurve zu beschreiben.

d) Es ist der Ort des bewegten Körpers zu den Zeiten $t = -3, -2, -1, 0, +1, +2, +3$ zu bestimmen. Diese Punkte sind auf der Bahnkurve zu markieren.

15. Lösen Sie die in Aufgabe 14 gestellte Aufgabe, wenn

für $x = g(t)$ gilt: $g(2) = 4$, $g'(2) = -2$, $g''(2) = -2$;

für $y = h(t)$ gilt: $h(2) = 1$, $h'(2) = +2$, $h''(2) = +2$!

16. Es ist durch eine Restgliedabschätzung die Richtigkeit der Formel $V_t = V_0(1 + 3\alpha t)$ für die räumliche Ausdehnung bei Erwärmung eines festen Körpers nachzuweisen (vgl. Lehrbuch der Physik, 9. Schuljahr, Seite 140).
17. Es sind $\sqrt{1,2}$, $\sqrt{1,3}$, $\sqrt{0,98}$, $\sqrt{0,92}$ mit Hilfe einer Approximation 1. Ordnung zu berechnen.
18. Dieselben Wurzeln sind mit Hilfe einer Approximation 2. Ordnung zu bestimmen.
19. Approximieren Sie die folgenden Funktionen an der Stelle $x = 0$ durch Polynome 1. bzw. 2. Grades und geben Sie das jeweilige Restglied an!
- a) $y = \frac{1}{1+x}$, b) $y = \frac{1}{(1+x)^2}$.
20. Es sind die zu den Funktionen der Aufgabe 19 gehörenden Kurven und die Schmiegungsparabeln 1. und 2. Ordnung zu zeichnen.
21. Zeichnen Sie zu den Funktionen $y = \sqrt{1+x}$ und $y = \frac{1}{\sqrt{1+x}}$ die Kurven und die Schmiegungsparabeln 1. und 2. Ordnung im Punkte $x_0 = 0$.
22. Geben Sie an, für welches Intervall der x -Werte die Approximation 1. Ordnung im Punkte mit der Abszisse $x_0 = 0$ bei den nachstehenden Funktionen zulässig ist, wenn der Fehler $p\%$ nicht überschreiten soll ($p = 0,1\%$, 1% , 10%).
- a) $y = \frac{1}{1+x}$, b) $y = \sqrt{1+x}$, c) $y = \frac{1}{\sqrt{1+x}}$.
23. Es ist bekannt, daß $\lim_{x \rightarrow 0} \frac{\sin x}{x} = 1$ ist (Lehrbuch der Mathematik, 11. Schuljahr, Seite 11), das heißt, daß für kleine Werte von x die Funktion $f(x) = \sin x$ durch $\varphi(x) = x$ ersetzt werden kann.
- a) Bestätigen Sie die Taylorformel $\sin x = x - \frac{x^3}{3!} \cos \vartheta x$ ($0 < \vartheta < 1$)!
- b) Geben Sie für das Restglied R_1 eine Abschätzung seines absoluten Betrages!
- c) Für welche Werte von x ist die Approximation 1. Ordnung brauchbar, wenn der entstehende Fehler höchstens $p = 0,1\%$; 1% ; 10% betragen darf?
24. Die Taylorformel $\sin x = x - \frac{x^3}{3!} + \frac{x^5}{5!} \cos \vartheta x$, ($0 < \vartheta < 1$), liefert eine Approximation 4. Ordnung für $\sin x$.
- a) Es ist die Darstellung abzuleiten.
- b) Welches ist die Approximation 3. Ordnung und wie unterscheidet sie sich von der Approximation 4. Ordnung?
- c) Der absolute Betrag des Restgliedes R_3 ist abzuschätzen.
- d) Warum enthalten die approximierenden Funktionen für $\sin x$ nur ungerade Potenzen von x ?
- e) Es sind die Funktion $y = \sin x$ und ihre Schmiegungsparabeln 1. und 3. Ordnung zu zeichnen.
25. a) Stellen Sie die Taylorformel für die Funktion $y = \operatorname{tg} x$ an der Stelle $x_0 = 0$ mit dem Restglied 1. Ordnung auf! Approximieren Sie die Funktion durch ein Polynom 1. Ordnung!
- b) Schätzen Sie die Größe des Fehlers ab!
- c) Bestimmen Sie das zulässige Intervall, wenn der Fehler $p = 0,1\%$; 1% ; 10% nicht überschreiten soll!

II. Grundbegriffe der Reihenlehre

3. Zahlenfolgen

Im Lehrbuch der Mathematik für das 10. Schuljahr wurde der Begriff der Zahlenfolge eingeführt.

Eine Zahlenfolge heißt **endlich**, wenn sie aus einer endlichen Anzahl von Gliedern besteht. Sie heißt **unendlich**, wenn die Anzahl ihrer Glieder über jede Grenze wächst.

Kann man alle Glieder einer Zahlenfolge nach der gleichen Vorschrift bilden und kann man diese Vorschrift durch einen mathematischen Ausdruck darstellen, der die Berechnung des k -ten Gliedes in Abhängigkeit vom Index k unmittelbar ermöglicht, so heißt dieser Ausdruck **allgemeines Glied** der Folge. Für die Zahlenfolge

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

schreibt man kurz das Symbol $\{a_k\}$. In der Regel gibt man an, welche Werte k annehmen kann:

$$\{a_k\} \text{ mit } k = 1, 2, \dots, n \quad \text{oder} \quad \{a_k\} \text{ } (k = 1, 2, \dots, n).$$

Man nennt dann k den **Laufindex**. Im allgemeinen durchläuft er die Folge der natürlichen Zahlen; in diesem Falle gibt er zugleich die Stellung des Gliedes an. In vielen Fällen kann der Anfangswert des Laufindex k auch die Zahl 0 sein. Es ist dann die Folge $\{a_k\}$ mit $k = 1, 2, 3, \dots$ identisch mit der Folge $\{a_{k+1}\}$ mit $k = 0, 1, 2, \dots$

Beispiel:

Die Folge $\{k\}$ mit $k = 1, 2, 3, \dots$ ist identisch mit der Folge $\{k+1\}$ mit $k = 0, 1, 2, \dots$

Im folgenden soll k , wenn nicht anders angegeben, alle natürlichen Zahlen durchlaufen.

Eine Zahlenfolge

$$a_1, a_2, a_3, \dots, a_k, \dots, a_n, \dots$$

heißt **wachsend**, wenn für jedes Glied a_k die Beziehung

$$a_k \leq a_{k+1}$$

gilt. Sie heißt **fallend**, wenn für jedes Glied a_k die Beziehung

$$a_k \geq a_{k+1}$$

gilt. Wachsende Folgen und fallende Folgen bezeichnet man als **monotone Folgen**.

Beispiele:

für wachsende Zahlenfolgen:

a) $5, 9, 13, 27, 38;$

b) $\{4k - 1\} = 3, 7, 11, 15, \dots;$

c) $\{3 \cdot 2^k\} = 6, 12, 24, 48, \dots;$

d) $\{3(k - 6)\} = -15, -12, -9, -6, -3, 0, 3, \dots;$

für fallende Zahlenfolgen:

e) $+3, 1, -5, -13;$

f) $\{1 - 3k\} = -2, -5, -8, -11, \dots;$

g) $\{9 - 2k\} = 7, 5, 3, 1, -1, -3, -5, \dots;$

h) $\left\{\frac{1}{3 \cdot 5^k}\right\} = \frac{1}{15}, \frac{1}{75}, \frac{1}{375}, \dots;$

für nicht monotone Zahlenfolgen:

i) $+3, -5, +8, -1, +13, -3;$

k) $\{(-1)^{k-1} \cdot k\} = 1, -2, +3, -4, \dots;$

l) $\left\{1 + \frac{(-1)^k}{k}\right\} = 0, +\frac{3}{2}, +\frac{2}{3}, +\frac{5}{4}, +\frac{4}{5}, \dots;$

m) $\left\{\frac{k + (-1)^k}{k^2}\right\} = 0, +\frac{3}{4}, +\frac{2}{9}, +\frac{5}{16}, +\frac{4}{25}, \dots$

Eine Zahlenfolge heißt **absolut wachsend (absolut fallend)**, wenn die Absolutbeträge ihrer Glieder wachsen (fallen). Nach diesen Definitionen können nicht-monotone Folgen absolut wachsend oder absolut fallend sein oder keine dieser beiden Eigenschaften haben.

4. Der Grenzwert einer Zahlenfolge

Bei einigen der angeführten Beispiele (vgl. h, l und m) bemerken wir, daß mit wachsendem Index k die Differenz der Glieder ständig kleiner wird und einem Grenzwert zustrebt. So streben die Glieder der Folge $\left\{\frac{1}{k}\right\}$ gegen Null, die der Folge $\left\{\frac{k-1}{k}\right\}$ gegen 1. Deshalb definieren wir:

Eine (endliche) Zahl g heißt **Grenzwert der Folge $\{a_k\}$** , wenn sich zu jeder noch so kleinen positiven Zahl ε ein geeigneter Index ν angeben läßt, von dem an für alle weiteren Glieder a_k (das heißt also die Glieder a_k mit $k \geq \nu$) der Folge die Beziehung

$$|a_k - g| < \varepsilon$$

gilt. Man schreibt in einem solchen Falle $\lim_{k \rightarrow \infty} a_k = g$. Mit anderen Worten: Die Beziehung $|a_k - g| < \varepsilon$ für $k \geq \nu$ bedeutet, daß alle Glieder a_k der Folge von einem bestimmten Glied a_ν an in einer Umgebung von g liegen, die kleiner als ε ist. Das sind aber fast alle Glieder mit Ausnahme nur endlich vieler. Das ε -Intervall um g kann beliebig klein gemacht werden. Mit ε ändert sich wohl der Index ν , aber stets liegen nur endlich viele Glieder a_k außerhalb des ε -Intervalls. Diese Tatsache können wir uns auf der Zahlengeraden geometrisch veranschaulichen.

Eine Zahlenfolge kann nicht gegen zwei verschiedene Grenzwerte konvergieren. Strebte sie nämlich gegen zwei Grenzwerte g und g' , so würden die Beziehungen

$$|a_k - g| < \varepsilon \quad \text{für } k \geq \nu$$

und

$$|a_k - g'| < \varepsilon \quad \text{für } k \geq \nu'$$

gelten. Dann lägen sowohl in einer ε -Umgebung von g als auch in einer ε -Umgebung von g' unendlich viele Glieder a_k . Da $\varepsilon > 0$ beliebig klein gemacht werden kann, läßt es sich so bestimmen, daß kein Glied a_k , das zur Umgebung von g gehört, zugleich der Umgebung von g' angehört und umgekehrt. Dann liegen aber sowohl außerhalb der ε -Umgebung von g als auch außerhalb der ε -Umgebung von g' unendlich viele Glieder a_k der Folge. Das aber ist ein Widerspruch zur Definition des Grenzwertes.

Eine Folge, die einen Grenzwert hat, heißt **konvergent**, jede andere Folge heißt **divergent**.

Offensichtlich ist die Folge

$$\left\{ \frac{1}{2k-1} \right\} = 1, \frac{1}{3}, \frac{1}{5}, \frac{1}{7}, \frac{1}{9}, \dots$$

konvergent.

Der Grenzwert einer Folge kann zur Folge selbst gehören (vgl. Beispiel *m*), muß es aber nicht (vgl. Beispiel *l*).

Eine divergente Zahlenfolge heißt **bestimmt divergent**, wenn ihre Glieder a_k gegen $+\infty$ oder $-\infty$ streben. In jedem anderen Falle heißt die Folge **unbestimmt divergent**.

Die Bedeutung der konvergenten Zahlenfolgen liegt darin, daß man mit ihnen neue Zahlen, ihre Grenzwerte, bilden kann. Dieser Sachverhalt ist uns bereits bekannt. Zum Beispiel kann man die Zahl π so bestimmen, daß man die wachsende Folge der Verhältnisse $\frac{u_n}{d}$ (Umfang des einbeschriebenen Vielecks zum Durchmesser) bzw. die fallende Folge der Verhältnisse $\frac{U_n}{d}$ (Umfang des umbeschriebenen Vielecks zum Durchmesser) bildet und feststellt, daß sie beide gegen denselben Grenzwert konvergieren. Die gesuchte Zahl π wird damit als der gemeinsame Grenzwert beider Folgen definiert. Ebenso wird die Zahl e , die Basis der natürlichen Logarithmen, durch einen Grenzwert bestimmt. Diese neuen Zahlen sind Grenzwerte unendlicher Zahlenfolgen. Da sich häufig kein allgemeines Glied angeben läßt, muß man die Glieder nacheinander berechnen. Man muß nach einer endlichen Anzahl von Gliedern die Folge abbrechen und kann deshalb nur angenäherte Werte für die darzustellende Zahl erhalten, aber — und das ist wichtig — Näherungswerte von beliebiger Genauigkeit.

An der Konvergenz einer Zahlenfolge und an ihrem Grenzwert ändert sich nichts, wenn endlich viele Glieder der Folge abgeändert werden, zum Beispiel wenn der Folge endlich viele Glieder vorangestellt werden. Diese Tatsache zeigt deutlich, daß der Begriff der Konvergenz und der des Grenzwertes nicht allein von den einzelnen Gliedern der Folge, sondern vielmehr von der Gesamtheit der Glieder der Folge als Ganzes bestimmt wird.

5. Reihen

a) Endliche Reihen

Unter einer endlichen Reihe verstehen wir die additive Verknüpfung der Glieder einer endlichen Zahlenfolge. Sind

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

die Glieder einer solchen Folge, so ist

$$a_1 + a_2 + a_3 + \dots + a_k + \dots + a_n$$

eine aus den Gliedern dieser Folge gebildete Reihe. Man schreibt dafür auch $\sum_{k=1}^n a_k$.

Die Glieder der Folge heißen dann Glieder der Reihe. Die Summe der Reihe ist das Ergebnis der Addition aller ihrer Glieder.

b) Unendliche Reihen

Analog definieren wir den Begriff der unendlichen Reihe. Dabei muß zunächst erklärt werden, was man unter einer Summe mit unendlich vielen Summanden versteht. Es sei

$$a_1, a_2, a_3, \dots, a_k, \dots$$

eine unendliche Zahlenfolge. Wir bilden die Folge ihrer Teilsummen

$$s_1, s_2, s_3, \dots, s_k, \dots,$$

wobei

$$s_1 = a_1; s_2 = a_1 + a_2; s_3 = a_1 + a_2 + a_3; \quad ; s_k = a_1 + a_2 + \dots + a_k$$

ist.

Existiert $\lim_{k \rightarrow \infty} s_k$ und ist $\lim_{k \rightarrow \infty} s_k = s$, so nennt man s die Summe der unendlichen Reihe $\sum_{k=1}^{\infty} a_k$, und die Reihe heißt konvergent. Existiert $\lim_{k \rightarrow \infty} s_k$ nicht, so hat die unendliche Reihe keine Summe, und die Reihe wird divergent genannt.

Auch bei Reihen unterscheidet man entsprechend dem Verhalten der Teilsummenfolge bestimmte und unbestimmte Divergenz. Während bei den endlichen Reihen die Summenbildung durch Addition der einzelnen Glieder in jedem Falle (wenigstens theoretisch) möglich ist, trifft dies bei den unendlichen Reihen nicht zu.

Beispiel:

$$\sum_{k=1}^{\infty} u_k = \frac{1}{1 \cdot 3} + \frac{1}{3 \cdot 5} + \frac{1}{5 \cdot 7} + \dots + \frac{1}{(2k-1)(2k+1)} + \dots$$

Das allgemeine Glied läßt sich auch in der Form $a_k = \frac{1}{2} \left(\frac{1}{2k-1} - \frac{1}{2k+1} \right)$ schreiben. Die k -te Teilsumme dieser Reihe ist dann

$$\begin{aligned} s_k &= \frac{1}{1 \cdot 3} + \frac{1}{3 \cdot 5} + \frac{1}{5 \cdot 7} + \dots + \frac{1}{(2k-1)(2k+1)} \\ &= \frac{1}{2} \left\{ \left(\frac{1}{1} - \frac{1}{3} \right) + \left(\frac{1}{3} - \frac{1}{5} \right) + \left(\frac{1}{5} - \frac{1}{7} \right) + \dots + \left(\frac{1}{2k-1} - \frac{1}{2k+1} \right) \right\} \\ s_k &= \frac{1}{2} \left(1 - \frac{1}{2k+1} \right). \end{aligned}$$

Die Teilsummenfolge $\{s_k\}$ ist also konvergent, ihr Grenzwert ist $s = \lim_{k \rightarrow \infty} s_k = \frac{1}{2}$. Die vorgelegte Reihe ist also ebenfalls konvergent und hat den Wert (die Summe) $s = \frac{1}{2}$.

Die Bestimmung der Teilsummen und der Grenzübergang sind unerlässlich bei Reihenuntersuchungen.

Wir wollen dies an unserem Beispiel zeigen. Es galt:

$$\sum_{k=1}^{\infty} u_k = \frac{1}{1 \cdot 3} + \frac{1}{3 \cdot 5} + \frac{1}{5 \cdot 7} + \dots$$

Wir können die Glieder dieser Reihe so umformen, daß wir

$$\sum_{k=1}^{\infty} u_k = \left(1 - \frac{2}{3}\right) + \left(\frac{2}{3} - \frac{3}{5}\right) + \left(\frac{3}{5} - \frac{4}{7}\right) + \dots + \left(\frac{k}{2k-1} - \frac{k+1}{2k+1}\right) + \dots$$

erhalten. Lösen wir die Klammern auf, so erhalten wir die Reihe

$$1 - \frac{2}{3} + \frac{2}{3} - \frac{3}{5} + \frac{3}{5} - \frac{4}{7} + \dots$$

Es ist dann die Folge der Teilsummen dieser Reihe:

$$s_1 = 1; \quad s_2 = 1 - \frac{2}{3}; \quad s_3 = 1 - \frac{2}{3} + \frac{2}{3}; \quad s_4 = 1 - \frac{2}{3} + \frac{2}{3} - \frac{3}{5}; \quad \dots$$

Wir erkennen, daß die Glieder dieser Folge abwechselnd den Wert 1 und einen Wert zwischen 0 und 1 annehmen. Die Folge der Teilsummen ist also unbestimmt divergent. Da uns bekannt ist, daß die ursprünglich vorgelegte Reihe konvergiert, müssen wir folgern, daß es falsch ist, jedes als Summe von zwei Summanden dargestellte Glied einer Reihe $\sum_{k=1}^{\infty} u_k$ als zwei Glieder aufzufassen. Wir kehren deshalb zu der Reihe

$$\left(1 - \frac{2}{3}\right) + \left(\frac{2}{3} - \frac{3}{5}\right) + \left(\frac{3}{5} - \frac{4}{7}\right) + \dots + \left(\frac{k}{2k-1} - \frac{k+1}{2k+1}\right) + \dots$$

zurück und bilden die Folge ihrer Teilsummen:

$$\begin{aligned} s_1 &= \left(1 - \frac{2}{3}\right); \\ s_2 &= \left(1 - \frac{2}{3}\right) + \left(\frac{2}{3} - \frac{3}{5}\right); \\ s_3 &= \left(1 - \frac{2}{3}\right) + \left(\frac{2}{3} - \frac{3}{5}\right) + \left(\frac{3}{5} - \frac{4}{7}\right); \\ s_4 &= \left(1 - \frac{2}{3}\right) + \left(\frac{2}{3} - \frac{3}{5}\right) + \left(\frac{3}{5} - \frac{4}{7}\right) + \left(\frac{4}{7} - \frac{5}{9}\right); \end{aligned}$$

Für die k -te Teilsumme ergibt sich dann:

$$s_k = \left(1 - \frac{2}{3}\right) + \left(\frac{2}{3} - \frac{3}{5}\right) + \dots + \left(\frac{k}{2k-1} - \frac{k+1}{2k+1}\right) = 1 - \frac{k+1}{2k+1},$$

daraus folgt

$$s = \lim_{k \rightarrow \infty} s_k = \lim_{k \rightarrow \infty} \left(1 - \frac{k+1}{2k+1}\right) = 1 - \lim_{k \rightarrow \infty} \frac{k+1}{2k+1} = 1 - \frac{1}{2} = \frac{1}{2}.$$

Aufgaben

1. Aus dem Bildungsgesetz des allgemeinen Gliedes a_k sind die ersten fünf Glieder der Folge $\{a_k\}$ mit $k = 1, 2, 3, \dots$ zu bestimmen.

$$\begin{array}{llll} \text{a)} a_k = \frac{k}{k+1} & \text{b)} a_k = \frac{k-1}{k} & \text{c)} a_k = \frac{1-k}{k} & \text{d)} a_k = (-1)^k \cdot k \\ \text{e)} a_k = \frac{1}{(-2)^k} & \text{f)} a_k = \frac{1}{k^2} & \text{g)} a_k = k! & \text{h)} a_k = \sqrt{k} \\ \text{i)} a_k = \frac{k}{2k+1} & \text{k)} a_k = \frac{3k-1}{1-2k} & \text{l)} a_k = e^{\frac{1}{k}} & \text{m)} a_k = \sin\left((k-1) \frac{\pi}{2}\right). \end{array}$$

2. Bestimmen Sie aus den gegebenen Gliedern der nachstehenden Folgen ein allgemeines Glied a_k ($k = 1, 2, 3, \dots$)!

$$\begin{array}{ll} \text{a)} 1, 8, 27, 64, 125, \dots & \text{b)} -3, -1, +1, +3, +5, \dots \\ \text{c)} 2, \frac{3}{2}, \frac{4}{3}, \frac{5}{4}, \frac{6}{5}, \dots & \text{d)} 1, -\frac{1}{3}, +\frac{1}{5}, -\frac{1}{7}, +\frac{1}{9}, \dots \\ \text{e)} \sqrt{a}; \sqrt[4]{a}; \sqrt[5]{a}; \sqrt[6]{a}; \sqrt[7]{a}; \dots & \text{f)} 9, 4, -1, -6, -11, \dots \end{array}$$

3. Es ist anzugeben, welche der Folgen in den Aufgaben 1 und 2

a) wachsen, b) fallen, c) nicht monoton sind!

4. Fertigen Sie eine graphische Darstellung der Folgen in den Aufgaben 1 und 2 an, indem Sie als unabhängige Variable den Index k , als abhängige Variable das Glied a_k wählen!

5. Es ist das Bildungsgesetz der Folgen in den Aufgaben 1 und 2 vom Index k ($k = 1, 2, 3, \dots$) auf den Index k' ($k' = 0, 1, 2, \dots$) umzuschreiben.

Welche geometrische Bedeutung hat diese Umschreibung?

6. Lassen Sie in den Bildungsgesetzen der Folgen in der Aufgabe 1 den Index k statt der natürlichen Zahlen die nicht negativen, ganzen Zahlen durchlaufen (ohne das Bildungsgesetz zu ändern)! Welchen Einfluß hat diese Änderung auf die Folge und ihre eventuelle Konvergenz?

7. Von den Folgen der Aufgaben 1 und 2 sind zu bestimmen

a) die konvergenten Folgen und ihre Grenzwerte,
b) die bestimmt divergenten Folgen,
c) die unbestimmt divergenten Folgen.

8. Geben Sie für die konvergenten Folgen der Aufgabe 1 an, von welchem Index k an das Glied a_k eine hinreichende Approximation des Grenzwertes darstellt, wenn die verlangte Genauigkeit für den Grenzwert mindestens

a) $\epsilon = 10^{-1}$, b) $\epsilon = 10^{-2}$, c) $\epsilon = 10^{-3}$ sein soll!

Anleitung:

Da der Grenzwert bereits bekannt ist, läßt sich aus der Beziehung $|g - a_k| \leq \epsilon$ bei bekanntem Bildungsgesetz der gesuchte Index k_0 bestimmen.

9. Für die Folgen der vorigen Aufgabe ist allgemein der Grenzwert k_0 bei vorgegebener Genauigkeitsschranke als Funktion von ε , das heißt $k_0 = k_0(\varepsilon)$, zu bestimmen.

10. Die ersten 5 Glieder der Reihen

$$\text{a) } \sum_{k=1}^{\infty} \frac{1}{k^2}, \quad \text{b) } \sum_{k=1}^{\infty} \frac{1}{k!}, \quad \text{c) } \sum_{k=1}^{\infty} \frac{1}{k}, \quad \text{d) } \sum_{k=1}^{\infty} \frac{k}{4^{k-1}}$$

sind anzugeben.

11. Stellen Sie fest, welche Glieder der Reihen in Aufgabe 10 keinen Beitrag mehr zur 4. Dezimale der Reihensumme liefern!

12. Es sind für die Reihen

$$\text{a) } 1 - \sum_{k=1}^{\infty} \frac{1}{k(k+1)} \quad (k = 1, 2, 3, \dots), \quad \text{b) } 1 + \sum_{k=1}^{\infty} \frac{1}{k(k+1)} \quad (k = 1, 2, 3, \dots)$$

das allgemeine Glied s_k der Teilsummenfolge zu bestimmen und daraus der Wert der Reihe herzuleiten.

13. Es sind die k -te Teilsumme der Reihe

$$1 + \sum_{k=1}^{\infty} (-1)^k \frac{2k+1}{k(k+1)} \quad (k = 1, 2, 3, \dots)$$

und ihre Summe zu bestimmen.

14. Es sind die k -te Teilsumme der Reihe

$$\sum_{k=1}^{\infty} (-1)^{k+1} \frac{2k+1}{k(k+1)} \quad \text{mit } k = 1, 2, 3, \dots$$

und ihre Summe zu bestimmen.

15. Stellen Sie die Teilsummenfolgen der Reihen in den Aufgaben 12 bis 14 in Abhängigkeit von ihrem Index graphisch dar!

16. Es sind die Glieder der Teilsummenfolgen der Reihen in den Aufgaben 12 bis 14 graphisch darzustellen.

17. Bestimmen Sie, von welchem Index k_0 ab die Glieder der Reihen in den Aufgaben 12 bis 14 keinen Beitrag mehr liefern zur

$$\text{a) } 2. \text{ Dezimale}, \quad \text{b) } 3. \text{ Dezimale}, \quad \text{c) } 4. \text{ Dezimale}$$

des Wertes der Reihe!

18. Für die Reihen in den Aufgaben 12 bis 14 ist festzustellen; von welchem Index k_0 ab die k -te Teilsumme der Reihen bis auf

$$\text{a) } 2 \text{ Dezimalen}, \quad \text{b) } 3 \text{ Dezimalen}, \quad \text{c) } 4 \text{ Dezimalen}$$

mit dem Grenzwert der betreffenden Reihen übereinstimmt!

Die Ergebnisse dieser Aufgabe sind mit den Ergebnissen der Aufgabe 17 zu vergleichen

III. Konvergenzuntersuchungen

6. Einfache Sätze über konvergente Reihen

Satz 5:

Wenn die Reihe $\sum_{k=1}^{\infty} u_k$ konvergiert, so konvergiert auch die Reihe

$$u_0 + \sum_{k=1}^{\infty} u_k = \sum_{k=0}^{\infty} u_k.$$

Beweis:

Es sei $\sum_{k=1}^{\infty} u_k = s$ und s_k die k -te Teilsumme der Reihe $\sum_{k=1}^{\infty} u_k$. Für die k -te Teilsumme s'_k der Reihe $u_0 + \sum_{k=1}^{\infty} u_k$ gilt dann $s'_k = u_0 + s_{k-1}$. Da $\lim_{k \rightarrow \infty} s_{k-1} = s$ ist, gilt

$$\lim_{k \rightarrow \infty} s'_k = \lim_{k \rightarrow \infty} u_0 + \lim_{k \rightarrow \infty} s_{k-1} = u_0 + \lim_{k \rightarrow \infty} s_{k-1} = u_0 + s = s'.$$

Diesen Satz kann man mehrmals hintereinander anwenden. Daraus folgt der *Satz 6:*

Die Konvergenz einer Reihe bleibt erhalten, wenn man eine endliche Anzahl von Gliedern hinzufügt oder wegnimmt.

Ferner gilt der

Satz 7:

Wenn die Reihe $\sum_{k=1}^{\infty} u_k$ konvergent ist, so ist stets $\lim_{k \rightarrow \infty} u_k = 0$.

Beweis:

Es sei $\lim_{k \rightarrow \infty} s_k = s$. Dann ist auch $\lim_{k \rightarrow \infty} s_{k-1} = s$. Nun besteht die Beziehung

$$u_k = s_k - s_{k-1}.$$

Daraus folgt

$$\lim_{k \rightarrow \infty} u_k = \lim_{k \rightarrow \infty} s_k - \lim_{k \rightarrow \infty} s_{k-1} = 0.$$

Die Gültigkeit der Beziehung $\lim_{k \rightarrow \infty} u_k = 0$ ist für die Konvergenz der Reihe $\sum_{k=1}^{\infty} u_k$ notwendig. Daß sie nicht hinreichend ist, zeigt das folgende Beispiel:

Für die Reihe

$$\frac{1}{\sqrt{1}} + \frac{1}{\sqrt{2}} + \frac{1}{\sqrt{3}} + \cdots + \frac{1}{\sqrt{k}} + \cdots$$

gilt

$$\lim_{k \rightarrow \infty} u_k = \lim_{k \rightarrow \infty} \frac{1}{\sqrt{k}} = 0.$$

Es ist jedoch

$$s_k = 1 + \frac{1}{\sqrt{2}} + \cdots + \frac{1}{\sqrt{k}}.$$

Da $1 > \frac{1}{\sqrt{2}} > \frac{1}{\sqrt{3}} > \frac{1}{\sqrt{4}} > \cdots > \frac{1}{\sqrt{k}}$ ist, besteht gewiß die Ungleichung

$$1 + \frac{1}{\sqrt{2}} + \frac{1}{\sqrt{3}} + \cdots + \frac{1}{\sqrt{k}} > \frac{1}{\sqrt{k}} + \frac{1}{\sqrt{k}} + \frac{1}{\sqrt{k}} + \cdots + \frac{1}{\sqrt{k}}.$$

Mithin gilt

$$s_k > k \cdot \frac{1}{\sqrt{k}} = \sqrt{k}.$$

Es existiert also kein Grenzwert der Teilsummenfolge. Daraus folgt, daß die Reihe divergent ist, obwohl die Beziehung $\lim_{n \rightarrow \infty} u_n = 0$ gilt.

Ist also bei einer vorgelegten Reihe diese Beziehung gültig, so kann man daraus noch nicht auf Konvergenz oder Divergenz schließen. Ist sie dagegen nicht gültig, so ist die Reihe sicher divergent.

Im folgenden wollen wir noch beweisen:

Satz 8:

Wenn $\sum_{k=1}^{\infty} u_k$ und $\sum_{k=1}^{\infty} v_k$ zwei Reihen mit nur positiven Gliedern sind und stets $u_k \leq v_k$ ist, so folgt aus der Konvergenz der Reihe $\sum_{k=1}^{\infty} v_k$ die Konvergenz der Reihe $\sum_{k=1}^{\infty} u_k$.

Beweis:

Nach der Voraussetzung sind alle u_k positiv. Ferner kann kein u_k größer als v_k werden. Demzufolge gilt für die k -ten Teilsummen

$$0 < u_1 + u_2 + \cdots + u_k \leq v_1 + v_2 + \cdots + v_k.$$

Wegen der Beziehung $0 < u_k \leq v_k$ gilt

$$0 < s_{u_k} = u_1 + u_2 + \cdots + u_k \leq s_{v_k} = v_1 + v_2 + \cdots + v_k.$$

Da aber s_{u_k} und s_{v_k} mit wachsendem k monoton zunehmen und s_{v_k} wegen der vorausgesetzten Konvergenz einem Grenzwert zustrebt, ergibt sich

$$0 < s_{u_k} \leq s_{v_k} < \lim_{k \rightarrow \infty} s_{v_k} = s.$$

Also nähert sich auch s_{u_k} mit zunehmendem k wachsend einem bestimmten Grenzwert $s' \leq s$. Die Reihe $\sum_{k=1}^{\infty} u_k$ konvergiert also.

Die bei diesen Konvergenzuntersuchungen angewendete Methode nennt man die **Methode der Reihenvergleichung**. Die Reihe $\sum_{k=1}^{\infty} v_k$ heißt **Majorante** der Reihe $\sum_{k=1}^{\infty} u_k$. Entsprechend gilt der

Satz 9:

Wenn $\sum_{k=1}^{\infty} u_k$ und $\sum_{k=1}^{\infty} v_k$ zwei Reihen mit nur positiven Gliedern sind und stets $u_k \geq v_k$ ist, so folgt aus der Divergenz der Reihe $\sum_{k=1}^{\infty} v_k$ die Divergenz der Reihe $\sum_{k=1}^{\infty} u_k$.

Der Beweis kann entsprechend geführt werden. In diesem Falle ist $\sum_{k=1}^{\infty} v_k$ Minorante der Reihe $\sum_{k=1}^{\infty} u_k$.

7. Konvergenzkriterien

Wir wollen nun hinreichende Bedingungen für die Konvergenz einer Reihe herleiten. Derartige Bedingungen nennt man Konvergenzkriterien.

a) Quotientenkriterium

Wenn für eine Reihe $\sum_{k=1}^{\infty} u_k$ mit nur positiven Gliedern von einem bestimmten $k = v$ an für alle weiteren Glieder die Bedingung

$$\frac{u_{k+1}}{u_k} \leq q < 1$$

erfüllt ist, so ist die Reihe konvergent.

Beweis:

Es ist

$$\frac{u_{v+1}}{u_v} \leq q; \quad \frac{u_{v+2}}{u_{v+1}} \leq q; \quad \frac{u_{v+3}}{u_{v+2}} \leq q; \quad \dots; \quad \frac{u_{v+n}}{u_{v+n-1}} \leq q;$$

$$u_{v+1} \leq q u_v; \quad u_{v+2} \leq q u_{v+1}; \quad u_{v+3} \leq q u_{v+2}; \quad \dots; \quad u_{v+n} \leq q u_{v+n-1}; \quad \dots$$

$$u_{v+1} \leq q u_v; \quad u_{v+2} \leq q^2 u_v; \quad u_{v+3} \leq q^3 u_v; \quad \dots; \quad u_{v+n} \leq q^n u_v;$$

Damit haben wir zu der Reihe

$$u_v + u_{v+1} + u_{v+2} + \dots + u_{v+n} + \dots$$

eine Majorante

$$u_v + u_v q + u_v q^2 + \dots + u_v q^n + \dots$$

gebildet, die eine unendliche geometrische Reihe mit $|q| < 1$ ist. Sie konvergiert, also konvergiert nach Satz 8 auch die Reihe

$$u_v + u_{v+1} + u_{v+2} + \dots + u_{v+n} + \dots$$

Die vorgelegte Reihe $\sum_{k=1}^{\infty} u_k$ unterscheidet sich von dieser Reihe nur dadurch, daß sie eine endliche Anzahl Glieder mehr enthält. Nach Satz 6 ist demnach auch die Reihe $\sum_{k=1}^{\infty} u_k$ konvergent. Damit ist die Behauptung bewiesen.

Die Bedingung $\frac{u_{n+1}}{u_n} \leq q < 1$ ist für die Konvergenz einer Reihe mit positiven Gliedern gewiß dann erfüllt, wenn

$$q' = \lim_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} < q < 1$$

gilt.

Beweis: Setzt man $\varepsilon = \frac{q - q'}{2}$, so ist für fast alle $\frac{u_{n+1}}{u_n}$ die Beziehung

$$\left| \frac{u_{n+1}}{u_n} - q' \right| < \varepsilon$$

erfüllt. Das heißt, es gilt für fast alle $\frac{u_{n+1}}{u_n}$ die Ungleichung

$$q' - \varepsilon < \frac{u_{n+1}}{u_n} < q' + \varepsilon.$$

Aus der Beziehung $\varepsilon = \frac{q - q'}{2}$ ergibt sich $2\varepsilon = q - q'$ oder $q' + \varepsilon - q - \varepsilon < q$. Also ist $q' + \varepsilon$ gewiß kleiner als q , mithin gilt

$$q' - \varepsilon < \frac{u_{n+1}}{u_n} < q' + \varepsilon < q.$$

Das bedeutet aber, daß für fast alle $\frac{u_{n+1}}{u_n}$ die Ungleichung

$$\frac{u_{n+1}}{u_n} \leq q < 1$$

erfüllt ist.

Es gilt also:

Eine Reihe $\sum_{k=1}^{\infty} u_k$ konvergiert, wenn es eine Zahl q gibt, für die die Beziehung

$$\lim_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} \leq q < 1$$

gilt.

Ferner gilt:

Eine Reihe $\sum_{k=1}^{\infty} u_k$ divergiert, wenn

$$\lim_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} = q > 1$$

ist.

Bei einer solchen Reihe können nämlich die Glieder mit wachsendem n nicht gegen Null streben; die Reihe muß also divergieren.

Dieses Kriterium, das als **Quotientenkriterium** bezeichnet wird, ist für die Konvergenz einer Reihe zwar hinreichend, aber nicht notwendig. Es kann also konvergente Reihen geben, für die dieses Kriterium nicht erfüllt ist.

So kann man zum Beispiel keine allgemeingültige Aussage mehr machen, wenn

$$\lim_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} = 1$$

ist. Reihen mit dieser Eigenschaft können konvergent oder divergent sein.

Daß das Quotientenkriterium für die Konvergenz einer Reihe nur hinreichend ist, zeigt das Beispiel

$$\sum_{k=1}^{\infty} \frac{1}{(2k+1)(2k-1)}.$$

Wir wissen bereits, daß diese Reihe konvergiert, trotzdem ist für sie das Quotientenkriterium nicht erfüllt. Es gilt nämlich

$$\begin{aligned} \lim_{k \rightarrow \infty} \frac{u_{k+1}}{u_k} &= \lim_{k \rightarrow \infty} \frac{(2k+1)(2k-1)}{(2k+3)(2k+1)} = \lim_{k \rightarrow \infty} \frac{2k-1}{2k+3} \\ &= \lim_{k \rightarrow \infty} \left(1 - \frac{4}{2k+3}\right) = 1. \end{aligned}$$

Für die Fälle, in denen dieses Kriterium versagt, gibt es schärfere Kriterien, die wir aber nicht behandeln wollen.

b) Kriterium von Leibniz

Wechselt bei einer Reihe von Glied zu Glied das Vorzeichen, so heißt die Reihe **alternierend**.

Wir wollen die alternierende Reihe

$$\sum_{k=1}^{\infty} (-1)^{k+1} u_k \quad (u_k > 0)$$

untersuchen. Wir setzen voraus, daß die Glieder u_k der Folge monoton gegen Null streben. Wir erkennen aus

$$\begin{aligned} s_1 &= u_1, \\ s_3 &= u_1 - (u_2 - u_3), \\ s_5 &= u_1 - (u_2 - u_3) + (u_4 - u_5), \\ &\vdots \end{aligned}$$

daß

$$s_1, s_3, s_5, \dots = \{s_{2k-1}\}$$

eine monoton fallende Folge ist, und aus

$$\begin{aligned} s_2 &= (u_1 - u_2), \\ s_4 &= (u_1 - u_2) + (u_3 - u_4), \\ s_6 &= (u_1 - u_2) + (u_3 - u_4) + (u_5 - u_6), \\ &\vdots \end{aligned}$$

daß

$$s_2, s_4, s_6, \dots = \{s_{2k}\}$$

eine monoton wachsende Folge ist. Andererseits ist $s_{2k-1} - s_{2k} = u_{2k} > 0$. Hieraus ergibt sich, daß alle s_{2k-1} größer als alle s_{2k} sind. Sie bilden eine monoton fallende Folge, die eine bestimmte Zahl k nicht unterschreitet. Eine solche Folge besitzt, wie wir hier nicht zeigen wollen, immer einen Grenzwert. Wir wollen ihn mit s' bezeichnen. Auch für die monoton wachsende Folge der s_{2k} kann man nachweisen, daß sie gegen einen Grenzwert s'' konvergieren.

Nun ergibt sich wegen der Voraussetzung $\lim_{k \rightarrow \infty} u_k = 0$ die Beziehung

$$\lim_{k \rightarrow \infty} u_k = \lim_{k \rightarrow \infty} (s_{2k-1} - s_{2k}) = \lim_{k \rightarrow \infty} s_{2k-1} - \lim_{k \rightarrow \infty} s_{2k} = s' - s'' = 0.$$

Es gilt also $s' = s''$.

Diese Aussage folgt aus der Tatsache, daß die Glieder u_k der Reihe eine monoton gegen Null strebende Folge bilden. Aus diesen Zusammenfassungen erkennen wir, daß die Teilsummenfolgen $s_{k+1}, s_{k+1}, \dots$ und $s_k, s_{k+2}, \dots$ von verschiedenen Seiten her gegen ein und denselben Grenzwert streben.

Es gilt also:

Eine alternierende Reihe $\sum_{k=1}^{\infty} (-1)^{k+1} u_k$ ist konvergent, wenn die Glieder u_k der Reihe monoton gegen Null streben.

Dieses Kriterium heißt **Kriterium von Leibniz**.

8. Absolute und bedingte Konvergenz

Wir wollen die Reihen $\sum_{k=0}^{\infty} x^k$ und $\sum_{k=0}^{\infty} (-x)^k$ für $0 < x < 1$ untersuchen. Auf die Reihe $\sum_{k=0}^{\infty} x^k$ wenden wir zunächst das Quotientenkriterium an: Es ist stets

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

und mithin

$$\lim_{n \rightarrow \infty} \frac{x^{n+1}}{x^n} = \lim_{n \rightarrow \infty} x = x.$$

Da $0 < x < 1$ vorausgesetzt wurde, konvergiert die Reihe. Die Reihe $\sum_{k=0}^{\infty} (-x)^k$ prüfen wir nach dem Kriterium von Leibniz: Offensichtlich strebt die Folge der Glieder x^k wegen $|x| < 1$ gegen Null. Also ist die Reihe konvergent.

Die beiden Reihen $\sum_{k=0}^{\infty} (-1)^k x^k$ und $\sum_{k=0}^{\infty} x^k$ unterscheiden sich nur durch die Vorzeichen der Glieder. Sie sind jedoch beide konvergent.

Definition:

Eine konvergente Reihe

$$u_1 + u_2 + u_3 + \dots$$

heißt absolut konvergent, wenn auch die Reihe

$$|u_1| + |u_2| + |u_3| + \dots$$

konvergiert.

Nach dieser Definition ist die Reihe $\sum_{k=0}^{\infty} (-x)^k$ für $0 < x < 1$ absolut konvergent. Nunmehr wollen wir die beiden Reihen $\sum_{k=1}^{\infty} (-1)^{k+1} \frac{1}{k}$ und $\sum_{k=1}^{\infty} \frac{1}{k}$ auf ihre Konvergenz hin untersuchen.

Die Reihe $\sum_{k=1}^{\infty} (-1)^{k+1} \frac{1}{k}$ konvergiert nach dem Kriterium von Leibniz.

Für die Reihe $\sum_{k=1}^{\infty} \frac{1}{k}$ gilt

$$\lim_{n \rightarrow \infty} \frac{u_{n+1}}{u_n} = \lim_{n \rightarrow \infty} \frac{\frac{1}{n+1}}{\frac{1}{n}} = 1,$$

also ist eine Entscheidung nach dem Quotientenkriterium nicht möglich. Wir müssen deshalb auf eine andere Weise entscheiden, ob die Reihe konvergiert oder divergiert. Zu diesem Zweck konstruieren wir eine Vergleichsreihe $\sum_{k=1}^{\infty} v_k$.

$$\begin{aligned} \sum_{k=1}^{\infty} \frac{1}{k} &= 1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \frac{1}{5} + \dots + \frac{1}{8} + \frac{1}{9} + \dots \\ &\quad + \frac{1}{16} + \frac{1}{17} + \dots + \frac{1}{32} + \frac{1}{33} + \dots \\ \sum_{k=1}^{\infty} v_k &= \frac{1}{2} + \frac{1}{2} + \frac{1}{4} + \frac{1}{4} + \frac{1}{8} + \dots + \frac{1}{8} + \frac{1}{16} + \dots \\ &\quad + \frac{1}{16} + \frac{1}{32} + \dots + \frac{1}{32} + \frac{1}{64} + \dots \end{aligned}$$

Man erkennt, daß für jedes k die Bedingung $\frac{1}{k} \geq v_k$ erfüllt ist. Die Reihe $\sum_{k=1}^{\infty} v_k$ ist also Minorante zur Reihe $\sum_{k=1}^{\infty} \frac{1}{k}$. Durch geeignete Zusammenfassung der

Glieder bilden wir aus der Reihe $\sum_{k=1}^{\infty} v_k$ die Reihe

$$1 + \frac{1}{2} + \frac{1}{2} + \frac{1}{2} + \dots$$

Diese Reihe divergiert sicher; denn es ist $\lim_{n \rightarrow \infty} u_n = \frac{1}{2} \neq 0$. Damit divergiert aber nach Satz 9 auch die Reihe $\sum_{k=1}^{\infty} \frac{1}{k}$.

Definition:

Eine konvergente Reihe

$$u_1 + u_2 + u_3 + \dots$$

heißt bedingt konvergent, wenn die Reihe

$$|u_1| + |u_2| + |u_3| + \dots$$

divergiert.

Nach dieser Definition ist die Reihe $\sum_{k=1}^{\infty} (-1)^{k+1} \frac{1}{k}$ nur bedingt konvergent.

Die Unterscheidung von absolut und bedingt konvergenten Reihen ist nicht nur für die Systematik bedeutungsvoll. Man kann vielmehr zeigen, daß auf die absolut konvergenten Reihen dieselben arithmetischen Grundgesetze anwendbar sind wie auf endliche Summen (Polynome) und daß man infolgedessen ebenso mit ihnen rechnen kann wie mit Polynomen. Daß dies für bedingt konvergente Reihen jedoch nicht gilt, wollen wir an einem Beispiel zeigen:

Wir gehen von der bedingt konvergenten Reihe

$$1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \cdots = L$$

aus und multiplizieren die Reihe gliedweise mit dem Faktor $\frac{1}{2}$:

$$\frac{1}{2} - \frac{1}{4} + \frac{1}{6} - \frac{1}{8} + \cdots = \frac{1}{2} L.$$

Wir addieren die beiden Reihen, indem wir die Addition gliedweise vornehmen:

$$1 + \frac{1}{2} - \frac{1}{2} - \frac{1}{4} + \frac{1}{3} + \frac{1}{6} - \frac{1}{4} - \cdots = \frac{3}{2} L.$$

Ordnen wir diese Reihe so, daß gleiche Nenner nebeneinanderstehen, so erhalten wir

$$1 + \frac{1}{2} - \frac{1}{2} - \frac{1}{4} - \frac{1}{4} + \frac{1}{3} - \frac{1}{8} - \frac{1}{8} + \frac{1}{5} - \cdots = \frac{3}{2} L,$$

mithin

$$1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \cdots = \frac{3}{2} L.$$

Daraus würde $L = \frac{3}{2} L$ folgen, was offensichtlich ein Widerspruch ist.

Da die Aussagen, von denen wir ausgegangen sind, richtig waren, muß die Ursache dieses Widerspruches in einer fehlerhaften Verknüpfung liegen, das heißt also in einer falschen Anwendung der Additionsgesetze. Tatsächlich kann man das kommutative Gesetz der Addition auf bedingt konvergente Reihen nicht anwenden.

Nach diesen Untersuchungen geben wir den folgenden Satz an.

Satz 10:

Konvergiert die Reihe

$$|u_1| + |u_2| + |u_3| + \cdots,$$

so konvergiert auch die Reihe

$$u_1 + u_2 + u_3 + \cdots.$$

Beweis:

Es ist

$$0 \leq \frac{|u_n| + u_n}{2} \leq |u_n| \quad \text{und} \quad 0 \leq \frac{|u_n| - u_n}{2} \leq |u_n|.$$

Die beiden Reihen

$$v_1 + v_2 + v_3 + \cdots \quad \text{mit} \quad v_k = \frac{|u_k| + u_k}{2}$$

und

$$w_1 + w_2 + w_3 + \dots \quad \text{mit} \quad w_k = \frac{|u_k| - u_k}{2}$$

konvergieren, denn die $|u_k|$ bilden eine konvergente Majorante zu beiden Reihen. Da sie beide nur aus positiven Gliedern bestehen, sind sie absolut konvergent; wir können daher die zweite Reihe von der ersten subtrahieren und erhalten dadurch

$$\begin{aligned} v_1 - w_1 + v_2 - w_2 + v_3 - w_3 + \dots, \\ (v_1 - w_1) + (v_2 - w_2) + (v_3 - w_3) + \dots \end{aligned}$$

Auch diese Reihe ist konvergent; denn es existiert

$$\sum_{k=1}^{\infty} (v_k - w_k) = \sum_{k=1}^{\infty} v_k - \sum_{k=1}^{\infty} w_k.$$

Wegen

$$v_k - w_k = \frac{|u_k| + u_k}{2} - \frac{|u_k| - u_k}{2} = u_k$$

konvergiert also auch die Reihe

$$u_1 + u_2 + u_3 + \dots$$

Damit ist der Satz bewiesen.

Durch diese Überlegungen erhalten wir unmittelbar den

Satz 11:

Die Reihe

$$u_1 + u_2 + u_3 + \dots$$

konvergiert absolut, wenn

$$\left| \frac{u_{k+1}}{u_k} \right| = \left| \frac{u_{k+1}}{u_k} \right| \leq q < 1$$

ist.

Aufgaben

1. Die folgenden Reihen sind auf Konvergenz hin zu untersuchen.

a) $\sum_{n=1}^{\infty} \frac{1}{n^2}$

d) $\sum_{n=1}^{\infty} \frac{1}{2^n}$

g) $\sum_{n=1}^{\infty} \frac{1}{n(1+n)^2}$

b) $\sum_{n=1}^{\infty} \frac{1}{\sqrt{n}}$

e) $\sum_{n=1}^{\infty} \frac{1}{1+n^2}$

h) $\sum_{n=1}^{\infty} \frac{1}{n(1+n)}$

c) $\sum_{n=1}^{\infty} \frac{n+1}{n}$

f) $\sum_{n=1}^{\infty} \frac{2^n}{n!}$

i) $\sum_{n=1}^{\infty} 2^n (\sin \alpha)^n$

2. Es ist zu untersuchen, für welche Werte von x die folgenden Reihen konvergieren.

a) $\sum_{n=1}^{\infty} x^n$

c) $\sum_{n=1}^{\infty} \frac{1}{x^n}$

e) $\sum_{n=1}^{\infty} n x^n$

b) $\sum_{n=1}^{\infty} \frac{x^n}{n}$

d) $\sum_{n=1}^{\infty} \frac{x^n}{n!}$

f) $\sum_{n=1}^{\infty} n^2 x^n$

IV. Potenzreihen

9. Allgemeines über Potenzreihen

Unter den Reihen haben die sogenannten Potenzreihen eine besondere Bedeutung. Potenzreihen sind Reihen von der Form

$$P(x) = a_0 + a_1 x + a_2 x^2 + \cdots + a_\mu x^\mu + \cdots = \sum_{k=0}^{\infty} a_k x^k.$$

Die Koeffizienten a_k sind konstante reelle Zahlen, während x eine Veränderliche ist. Bei der Konvergenzuntersuchung solcher Reihen ist für x jeweils ein fester Wert $x = \xi$ anzunehmen.

Satz 12:

Jede Potenzreihe $P(x)$ ist für $x=0$ konvergent.

Beweis:

Die k -te Teilsumme der Potenzreihe ist

$$s_k(x) = a_0 + a_1 x + a_2 x^2 + \cdots + a_k x^k.$$

Ist $x = 0$, so werden alle Glieder außer dem Glied a_0 gleich Null, und alle Teilsummen reduzieren sich auf den Wert a_0 . Also ist

$$P(0) = \lim_{k \rightarrow \infty} s_k(0) = a_0.$$

Satz 13:

Ist die Potenzreihe $P(x)$ für $x = \xi > 0$ konvergent, so ist sie für alle x zwischen $-\xi$ und $+\xi$, das heißt für $-\xi < x < +\xi$ konvergent (und zwar absolut konvergent).

Beweis:

Da die Reihe $P(\xi) = a_0 + a_1 \xi + a_2 \xi^2 + \cdots$ laut Voraussetzung konvergiert, so streben die Glieder $a_k \xi^k$ mit wachsendem k gegen Null, sie sind also von einem bestimmten $k = \nu$ an dem Betrage nach sämtlich kleiner als eine feste Zahl Z , das heißt

$$|a_k \xi^k| < Z \text{ für } k \geq \nu.$$

Das allgemeine Glied der Potenzreihe

$$P(x) = a_0 + a_1 x + a_2 x^2 + \cdots$$

läßt sich nun folgendermaßen abschätzen:

$$\left| a_k x^k \right| = \left| a_k \xi^k \right| \cdot \left| \frac{x}{\xi} \right|^k < Z \cdot \left| \frac{x}{\xi} \right|^k$$

Die geometrische Reihe $\sum_{k=\nu}^{\infty} Z \cdot \left| \frac{x}{\xi} \right|^k = \sum_{k=\nu}^{\infty} v_k$ ist eine Majorante der Reihe $\sum_{k=\nu}^{\infty} a_k x^k$.

Wenn $\left| \frac{x}{\xi} \right| < 1$ ist, so ist $\sum_{k=\nu}^{\infty} v_k$ eine konvergente geometrische Reihe, mithin

ist auch $\sum_{k=0}^{\infty} a_k x^k$ konvergent. Dann ist aber auch die Reihe $P(x) = \sum_{k=0}^{\infty} a_k x^k$ konvergent (und zwar absolut konvergent) für alle x , die der Bedingung $-\xi < x < +\xi$ genügen.

Existiert noch ein Wert $\xi_1 > \xi$, für den die Reihe konvergiert, so konvergiert sie auch im Intervall $-\xi_1 < x < +\xi_1$. Das größte Intervall, das auf solche Weise durch Erweiterung gefunden werden kann, heißt das **Konvergenzintervall** der Potenzreihe, seine halbe Länge der **Konvergenzradius** der Potenzreihe.

Aus dieser Herleitung folgt der

Satz 14:

Eine Potenzreihe konvergiert innerhalb ihres **Konvergenzintervalls**, außerhalb des **Konvergenzintervalls** divergiert sie.

Über das Verhalten an den Grenzen des **Konvergenzintervalls** kann man keine allgemeingültige Aussage machen. Man muß das in jedem einzelnen Fall untersuchen.

Wenn für eine Potenzreihe $\sum_{k=0}^{\infty} a_k x^k$ der Ausdruck $\lim_{k \rightarrow \infty} \left| \frac{a_k}{a_{k+1}} \right|$ existiert und a_k stets von Null verschieden ist, so hat der Konvergenzradius r dieser Reihe den Wert

$$r = \lim_{k \rightarrow \infty} \left| \frac{a_k}{a_{k+1}} \right|.$$

Wenn $r = 0$ ist, konvergiert die Potenzreihe nur für den Wert $x = 0$. Die Beweise für diese Behauptungen wollen wir nicht führen. Eine Potenzreihe, die nur für $x = 0$ konvergiert, heißt eine **nirgends konvergente** Potenzreihe. Eine Potenzreihe, für die $r = \lim_{k \rightarrow \infty} \left| \frac{a_k}{a_{k+1}} \right| = \infty$ gilt, ist für alle x konvergent. Man nennt sie eine **beständig konvergente** Potenzreihe.

Beispiel:

Gegeben ist die Potenzreihe

$$L(x) = x - \frac{x^2}{2} + \frac{x^3}{3} - \frac{x^4}{4} + \frac{x^5}{5} - + \dots$$

Für diese Reihe ist $\left| \frac{a_k}{a_{k+1}} \right| = \frac{k+1}{k}$, also ist $r = \lim_{k \rightarrow \infty} \frac{k+1}{k} = 1$. Die vorgelegte Potenzreihe ist also absolut konvergent für alle x mit $-1 < x < +1$. Für $x = 1$ erhalten wir die schon bekannte bedingt konvergente Reihe

$$L(1) = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \frac{1}{5} - + \dots,$$

für $x = -1$ ergibt sich die sogenannte harmonische Reihe

$$L(x) = (-1) \left(1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \dots \right),$$

die divergent ist. Die Potenzreihe ist also konvergent für $-1 < x \leq +1$.

Da die Potenzreihen innerhalb ihrer Konvergenzintervalle absolut konvergent sind, kann man mit diesen Reihen innerhalb ihrer Konvergenzintervalle wie mit Polynomen rechnen. Man kann sie daher als spezielle Funktionen mit einer unabhängigen Variablen innerhalb dieser Intervalle gliedweise differenzieren und integrieren.

10. Die Taylorsche Reihe

Wir betrachten die Taylorsche Formel für eine beliebige Funktion $f(x)$. Es gilt unter den genannten Voraussetzungen

$$f(x_0 + h) = f(x_0) + \frac{h}{1!} f'(x_0) + \frac{h^2}{2!} f''(x_0) + \dots + \frac{h^{n-1}}{(n-1)!} f^{(n-1)}(x_0) + R_n(x_0; h).$$

Wenn $f(x)$ Ableitungen beliebig hoher Ordnung hat und

$$\lim_{n \rightarrow \infty} R_n(x_0; h) = 0$$

im ganzen Intervall von x_0 bis $x_0 + h$ ist, gilt

$$f(x_0 + h) = f(x_0) + \frac{h}{1!} f'(x_0) + \frac{h^2}{2!} f''(x_0) + \dots.$$

Beweis:

Es ist

$$\begin{aligned} \lim_{n \rightarrow \infty} \left[f(x_0) + \frac{h}{1!} f'(x_0) + \frac{h^2}{2!} f''(x_0) + \dots + \frac{h^{n-1}}{(n-1)!} f^{(n-1)}(x_0) \right] \\ = \lim_{n \rightarrow \infty} [f(x_0 + h) - R_n] = \lim_{n \rightarrow \infty} f(x_0 + h) - \lim_{n \rightarrow \infty} R_n. \end{aligned}$$

Unter der Voraussetzung

$$\lim_{n \rightarrow \infty} R_n = 0$$

ergibt sich, da $f(x_0 + h)$ ein von n unabhängiger, wohlbestimmter Wert ist,

$$\lim [f(x_0 + h) - R_n] = f(x_0 + h).$$

Daraus folgt

$$f(x_0 + h) = f(x_0) + \frac{h}{1!} f'(x_0) + \frac{h^2}{2!} f''(x_0) + \dots.$$

Damit ist die Behauptung bewiesen.

Diese unendliche Reihe heißt **Taylorsche Reihe**.

Wir wollen jetzt untersuchen, unter welcher Bedingung das Restglied R_n gegen Null konvergiert. Für unsere Untersuchungen wählen wir wieder die Lagrangesche Form:

$$R_n = \frac{h^n}{n!} f^{(n)}(x_0 + \vartheta h) \quad (0 < \vartheta < 1).$$

Läßt sich für alle $n = 1, 2, 3, \dots$ und für jeden Wert von ϑ im Intervall $0 < \vartheta < 1$ eine Konstante c angeben mit der Bedingung

$$|f^{(n)}(x_0 + \vartheta h)| < c,$$

so konvergiert $R_n = \frac{h^n}{n!} f^{(n)}(x_0 + \vartheta h)$ gegen Null. Es ist dann nämlich

$$|R_n| < \frac{|h|^n}{n!} \cdot c.$$

Wenn $\lim_{n \rightarrow \infty} \frac{|h|^n}{n!} \cdot c = c \cdot \lim_{n \rightarrow \infty} \frac{|h|^n}{n!} = 0$ ist, gilt $\lim_{n \rightarrow \infty} |R_n| = 0$. Wir müssen also

$$\lim_{n \rightarrow \infty} \frac{|h|^n}{n!}$$

untersuchen. Dazu konstruieren wir uns eine passende Reihe, nämlich

$$1 + \frac{|h|}{1!} + \frac{|h|^2}{2!} + \frac{|h|^3}{3!} + \dots$$

Nach dem Quotientenkriterium ist

$$\frac{u_{n+1}}{u_n} = \frac{|h|^{n+1}}{(n+1)!} \cdot \frac{n!}{|h|^n} = \frac{|h|}{n+1}.$$

Für $n \rightarrow \infty$ geht dieser Ausdruck gegen Null. Die Reihe konvergiert also. Demnach muß auch nach Satz 7 $\lim_{n \rightarrow \infty} u_n = \lim_{n \rightarrow \infty} \frac{|h|^n}{n!} = 0$ sein. Es gilt also auch $c \cdot \lim_{n \rightarrow \infty} \frac{|h|^n}{n!} = 0$ und mithin

$$\lim_{n \rightarrow \infty} R_n = 0.$$

Damit erhalten wir den

Satz 15:

Das Restglied R_n der Taylorschen Formel zu einer Funktion $f(x)$ konvergiert gegen Null, wenn für alle n und für alle Werte von ϑ die Bedingung

$$|f^{(n)}(x_0 + \vartheta h)| < c$$

erfüllt ist. Die Funktion $f(x)$ ist dann in eine Taylorsche Reihe entwickelbar.

Diese Bedingung ist hinreichend, aber nicht notwendig.

Wird die Taylorsche Reihe einer Funktion $y = f(x)$ nach dem n -ten Glied abgebrochen, so erhalten wir eine Approximation n -ter Ordnung der Funktion $f(x)$. Der durch die Approximation entstandene Fehler wird durch das Restglied R_n abgeschätzt. Die den Näherungspolyomen zugeordneten Kurven heißen Schmiegungsparabeln der Funktion von der Ordnung $(n - 1)$.

Im folgenden werden wir für einige spezielle Funktionen untersuchen, ob sie in eine Taylorsche Reihe entwickelt werden können. Wir müssen dazu für die jeweils zu betrachtende Funktion $f(x)$ prüfen, ob unter den genannten Voraussetzungen die Beziehung $\lim R_n(x_0; h) = 0$ gilt und ob an der Stelle x_0 die Ableitungen beliebiger Ordnung existieren.

11. Die Exponentialreihe

Die Darstellung der Exponentialfunktion $y = e^x$ durch eine Reihe nach Potenzen von x ist besonders einfach, da die Ableitungen beliebiger Ordnung wieder die Funktion $y = e^x$ ergeben. Es ist stets $y^{(k)}(x) = e^x$ und mithin $y^{(k)}(0) = 1$ für jedes endliche k . Damit ergibt sich nach dem Satz von Taylor

$$y = e^x = 1 + \frac{1}{1!}x + \frac{1}{2!}x^2 + \frac{1}{3!}x^3 + \cdots + R_n.$$

Wir prüfen, ob das Restglied R_n gegen Null konvergiert. Nach Satz 15 ist das der Fall, wenn die Bedingung

$$|f^{(n)}(x_0 + \vartheta h)| < c$$

für alle $n = 1, 2, 3, \dots$ und alle Werte von ϑ erfüllt ist. Nun ist $f^{(n)}(x_0 + \vartheta h) = e^{x_0 + \vartheta h}$. Da für ϑ die Relation $0 < \vartheta < 1$ gilt, ist

$$e^{x_0} < e^{x_0 + \vartheta h} < e^{x_0 + h},$$

mithin gilt

$$|e^{x_0 + \vartheta h}| < c \quad \text{für } c = e^{x_0 + h}.$$

Wir können also die Exponentialfunktion $y = e^x$ als Taylorsche Reihe von der Form

$$y = e^x = 1 + \frac{1}{1!}x + \frac{1}{2!}x^2 + \cdots \quad (1)$$

darstellen. Wir prüfen noch mit dem Quotientenkriterium, für welche Werte von x die Reihe konvergiert. Es ist

$$\lim_{k \rightarrow \infty} \left| \frac{u_{k+1}}{u_k} \right| = \lim_{k \rightarrow \infty} \left| \frac{x^{k+1} \cdot k!}{(k+1)! \cdot x^k} \right| = \lim_{k \rightarrow \infty} \left| \frac{x}{k+1} \right| = 0$$

für jeden beliebigen, aber fest gewählten Wert von x . Die Reihe (1) ist also für alle x absolut konvergent.

Setzt man in dieser Reihe für x den Wert $x = 1$ ein, so erhält man eine exakte Darstellung der Zahl e :

$$e = 1 + \frac{1}{1!} + \frac{1}{2!} + \frac{1}{3!} + \frac{1}{4!} + \cdots \quad (2)$$

Bedenken wir, daß $y = a^x = e^{x \ln a}$ ist, so erhalten wir für die Funktion $y = a^x$ ($a > 0$) die folgende Potenzreihenentwicklung:

$$y = a^x = 1 + \frac{\ln a}{1!}x + \frac{(\ln a)^2}{2!}x^2 + \frac{(\ln a)^3}{3!}x^3 + \cdots.$$

Auch diese Reihe ist für alle x -Werte konvergent, da die Reihe (1) mit x auch für $x \cdot \ln a$ konvergiert.

12. Die Reihen für $\sin x$ und $\cos x$

Wir untersuchen die Ableitungen der Funktion $y = \sin x$. Ist die Ordnung k der Ableitung gerade, also $k = 2m$, so ist $y^{(k)} = (-1)^m \cdot \sin x$; ist dagegen k ungerade, also $k = 2m + 1$, so ist $y^{(k)} = (-1)^m \cdot \cos x$. Für $x = 0$ sind daher alle Ableitungen gerader Ordnung Null, während die Ableitungen der ungeraden Ordnung $k = 2m + 1$ den Wert $y^{(k)}(0) = (-1)^m$ haben. Daraus ergibt sich für die Funktion $y = \sin x$ die Entwicklung

$$\sin x = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \dots + (-1)^{n-1} \frac{x^{2n-1}}{(2n-1)!} + R_n.$$

Dabei ist

$$R_n = (-1)^n \frac{x^{2n+1}}{(2n+1)!} \cos \vartheta x,$$

das heißt $f^{(n)}(\vartheta x) = (-1)^n \cdot \cos \vartheta x$, da wir hier in dem Ausdruck $f^{(n)}(x_0 + \vartheta h)$ für x_0 den Wert 0 und für h den Wert x gewählt haben. Da $|\cos \vartheta x| \leq 1$ für alle $n = 1, 2, 3, \dots$ und alle ϑ zwischen 0 und 1 gilt, konvergiert R_n für $n \rightarrow \infty$ gegen Null. Die Funktion $y = \sin x$ ist also in eine Taylorsche Reihe entwickelbar:

$$y = \sin x = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \dots \quad (3)$$

Auch diese Reihe ist für alle x -Werte absolut konvergent. Ihr allgemeines Glied ist nämlich $u_k = (-1)^{k-1} \cdot \frac{x^{2k-1}}{(2k-1)!}$. Bilden wir den Quotienten $\frac{u_{k+1}}{u_k}$, so erhalten wir

$$\frac{u_{k+1}}{u_k} = \frac{(-1)^k \cdot x^{2k+1} \cdot (2k-1)!}{(-1)^{k-1} \cdot x^{2k-1} \cdot (2k+1)!} = -\frac{x^2}{2k(2k+1)}.$$

Daraus folgt, daß der Grenzwert dieses Quotienten mit wachsendem k für jeden beliebigen, aber fest gewählten Wert von x gegen Null strebt. Damit ist die Konvergenz der Reihe für alle Werte von x nachgewiesen.

In derselben Weise erhält man die für jeden endlichen x -Wert absolut konvergente Reihe

$$y = \cos x = 1 - \frac{x^2}{2!} + \frac{x^4}{4!} - \frac{x^6}{6!} + \dots \quad (4)$$

Die Entwicklung der $\cos$ -Reihe und die Restglieduntersuchung werden analog der $\sin$ -Reihe durchgeführt.

Aus der Reihe (3) folgt $\sin(-x) = -\sin x$, aus der Reihe (4) $\cos(-x) = \cos x$. Werden die Funktionen $y = \sin x$ und $y = \cos x$ durch die Reihen (3) und (4) definiert, so entnimmt man diesen Definitionen die Differentiationsregeln

$$\frac{d \sin x}{dx} = \cos x \quad \text{bzw.} \quad \frac{d \cos x}{dx} = -\sin x,$$

da absolut konvergente Reihen in ihrem Konvergenzbereich gliedweise differenziert werden können.

13. Die binomische Reihe

Im Lehrbuch der Mathematik für das 11. Schuljahr haben wir die binomische Formel

$$(a + b)^n = \sum_{v=0}^n \binom{n}{v} a^{n-v} b^v = \binom{n}{0} a^n + \binom{n}{1} a^{n-1} b + \dots + \binom{n}{n} b^n$$

kennengelernt. Wir haben ihre Gültigkeit für ganze positive Exponenten n nachgewiesen. Nunmehr wollen wir diese Einschränkung fallen lassen und die Funktion

$$y = (a + x)^\beta$$

mit beliebigen rationalen Exponenten β in eine Reihe entwickeln. Diese Funktion ist beliebig oft differenzierbar für alle $x \neq -a$. An der Stelle $x = -a$ sind die Funktion und ihre k -ten Ableitungen nicht definiert, sobald $\beta - k < 0$ ist, das heißt, die Bedingungen für die Entwicklung des Taylorschen Satzes sind dann nicht erfüllt. Es gilt

$$f^{(k)}(x) = \beta(\beta - 1)(\beta - 2) \cdots (\beta - [k - 1]) (a + x)^{\beta - k}.$$

Setzen wir $x = 0$, so erhalten wir

$$f^{(k)}(0) = \beta(\beta - 1)(\beta - 2) \cdots (\beta - [k - 1]) a^{\beta - k}.$$

Nach Multiplikation dieser Gleichung mit $\frac{x^k}{k!}$ ergibt sich

$$\frac{x^k}{k!} f^{(k)}(0) = \binom{\beta}{k} a^{\beta - k} x^k.$$

Wir erkennen, daß die linke Seite dieser Gleichung die uns bekannte Form des allgemeinen Gliedes des Taylorschen Lehrsatzes für $x_0 = 0$ und $h = x$ ist. Approximieren wir die Funktion $y = (a + x)^\beta$ durch ein Polynom $(m - 1)$ -ter Ordnung, so erhalten wir

$$y = (a + x)^\beta = a^\beta + \binom{\beta}{1} a^{\beta-1} x + \binom{\beta}{2} a^{\beta-2} x^2 + \cdots + \binom{\beta}{m-1} a^{\beta-m+1} x^{m-1} + R_m.$$

Das Restglied R_m nach Lagrange hat die Form

$$R_m = \frac{\beta(\beta - 1)(\beta - 2) \cdots (\beta - [m - 1])}{1 \cdot 2 \cdot 3 \cdots m} x^m (a + \vartheta x)^{\beta - m} \quad (0 < \vartheta < 1).$$

Der allgemeine Beweis, daß auch für diesen Fall $\lim_{m \rightarrow \infty} R_m = 0$ gilt, kann mit den zur Verfügung stehenden Mitteln nicht geführt werden, da der Ausdruck

$$f^{(m)}(\vartheta x) = \beta(\beta - 1)(\beta - 2) \cdots (\beta - [m - 1]) (a + \vartheta x)^{\beta - m}$$

mit $x_0 = 0$ und $h = x$ nicht ohne weiteres abgeschätzt werden kann.

Die Funktion $y = (a + x)^\beta$ läßt sich in die Reihe

$$(a + x)^\beta = a^\beta + \binom{\beta}{1} a^{\beta-1} x + \binom{\beta}{2} a^{\beta-2} x^2 + \binom{\beta}{3} a^{\beta-3} x^3 + \dots \quad (5)$$

entwickeln. Es ist noch zu prüfen, für welche Werte von x diese Reihe konvergiert.

Für positive ganze β bricht die binomische Reihe für jeden Wert von x wegen $\binom{\beta}{k} = 0$ für $k > \beta$ nach $k = \beta$ Gliedern ab. Sie ist also endlich und somit für jeden Wert von x konvergent.

Ist β nicht positiv ganz, so sind für $x \neq 0$ alle Glieder von Null verschieden und die Reihe ist unendlich.

Wir prüfen deshalb die Konvergenz mit Hilfe des Quotientenkriteriums. Es ist

$$u_k = \binom{\beta}{k-1} a^{\beta-(k-1)} x^{k-1}$$

und

$$u_{k+1} = \binom{\beta}{k} a^{\beta-k} x^k,$$

mithin

$$\frac{u_{k+1}}{u_k} = \frac{\binom{\beta}{k} a^{\beta-k} x^k}{\binom{\beta}{k-1} a^{\beta-(k-1)} x^{k-1}} = \frac{\binom{\beta}{k} x}{\binom{\beta}{k-1} a} = \frac{\beta - k + 1}{k} \cdot \frac{x}{a} = \left(\frac{\beta + 1}{k} - 1 \right) \frac{x}{a}.$$

Für $k \rightarrow \infty$ geht $\frac{\beta + 1}{k}$ gegen Null und demnach $\frac{u_{k+1}}{u_k}$ gegen $-\frac{x}{a}$. Daraus ergibt sich, daß die Reihe konvergiert, wenn

$$\left| \frac{x}{a} \right| < 1, \text{ das heißt } |x| < |a| \text{ ist.}$$

Beispiel:

Es ist $\sqrt{1,2}$ auf 6 Dezimalen genau zu berechnen. Wir gehen von der Funktion $y = \sqrt{1+x} = (1+x)^{\frac{1}{2}}$ aus und setzen dann $x = 0,2$. Es ist

$$(1+x)^{\frac{1}{2}} = 1 + \binom{\frac{1}{2}}{1} x + \binom{\frac{1}{2}}{2} x^2 + \binom{\frac{1}{2}}{3} x^3 + \binom{\frac{1}{2}}{4} x^4 + \dots$$

Nun ist

$$\begin{aligned} \binom{\frac{1}{2}}{k} &= \frac{\frac{1}{2} \cdot \left(-\frac{1}{2}\right) \cdot \left(-\frac{3}{2}\right) \cdot \left(-\frac{5}{2}\right) \cdots \left(-\frac{2k-3}{2}\right)}{1 \cdot 2 \cdot 3 \cdot 4 \cdots k} \\ &= (-1)^{k-1} \cdot \frac{1 \cdot 1 \cdot 3 \cdot 5 \cdot 7 \cdots (2k-3)}{2 \cdot 4 \cdot 6 \cdot 8 \cdot 10 \cdots 2k} \quad (k=2, 3, 4, \dots). \end{aligned}$$

Wir erhalten der Reihe nach für die Binomialkoeffizienten $\binom{1}{k}$ mit $k = 0, 1, 2, 3, \dots$ die Werte

$$1, \frac{1}{2}, -\frac{1}{8}, \frac{1}{16}, -\frac{5}{128}, \frac{7}{256}, -\frac{21}{1024}, \frac{33}{2048}, -\frac{429}{32768}, \dots$$

Da der Wurzelwert auf 6 Dezimalen berechnet werden soll, geben wir die

$$a_k = \binom{1}{k} 0,2^k \quad (k = 0, 1, 2, 3, \dots)$$

auf 8 Dezimalen an. Dann ist

$a_0 = 1,000\,000\,00$	$a_2 = -0,005\,000\,00$
$a_1 = 0,100\,000\,00$	$a_4 = -0,000\,062\,50$
$a_3 = 0,000\,500\,00$	$a_6 = -0,000\,001\,31$
$a_5 = 0,000\,008\,75$	$a_8 = -0,000\,000\,03$
$a_7 = 0,000\,000\,20$	
$\sqrt{1,2} = 1,100\,508\,95$	$-0,005\,063\,84 = 1,095\,445\,11$

Der auf 6 Stellen berechnete Wurzelwert ist also $\sqrt{1,2} = 1,095445$.

Durch Ausrechnen läßt sich nachprüfen, daß $1,095445^2$ um weniger als $0,3 \cdot 10^{-6}$ von $1,2$ abweicht!

Die binomische Reihe konvergiert um so besser, je kleiner der absolute Betrag des für x zu wählenden Wertes ist. Diese Tatsache macht man sich bei der Berechnung von Wurzelwerten zunutze, indem man den Radikanden entsprechend umformt.

So läßt sich zum Beispiel $\sqrt{2}$ wie folgt umformen:

$$\text{a) } \sqrt{2} = \sqrt{4-2} = \sqrt{4\left(1-\frac{1}{2}\right)} = 2\left(1-\frac{1}{2}\right)^{\frac{1}{2}}$$

oder

$$\text{b) } \sqrt{2} = \sqrt{\frac{49}{25} \cdot \frac{50}{49}} = \frac{7}{5} : \sqrt{\frac{49}{50}} = \frac{7}{5} \left(1 - \frac{1}{50}\right)^{-\frac{1}{2}}.$$

In der Abbildung 10 sind die zu den approximierenden Funktionen gehörigen Schmiegungsparabeln 0-ter bis 4-ter Ordnung der Reihe $\sum_{k=0}^{\infty} x^k = \frac{1}{1-x}$ dargestellt.

Die Abbildung zeigt, daß diese Reihe weder bei $x = +1$ (bestimmte Divergenz) noch bei $x = -1$ (unbestimmte Divergenz) konvergiert, während sie für alle x zwischen $+1$ und -1 konvergiert, und zwar um so besser, je näher der x -Wert bei Null liegt und je höher die Ordnung der Approximation ist.

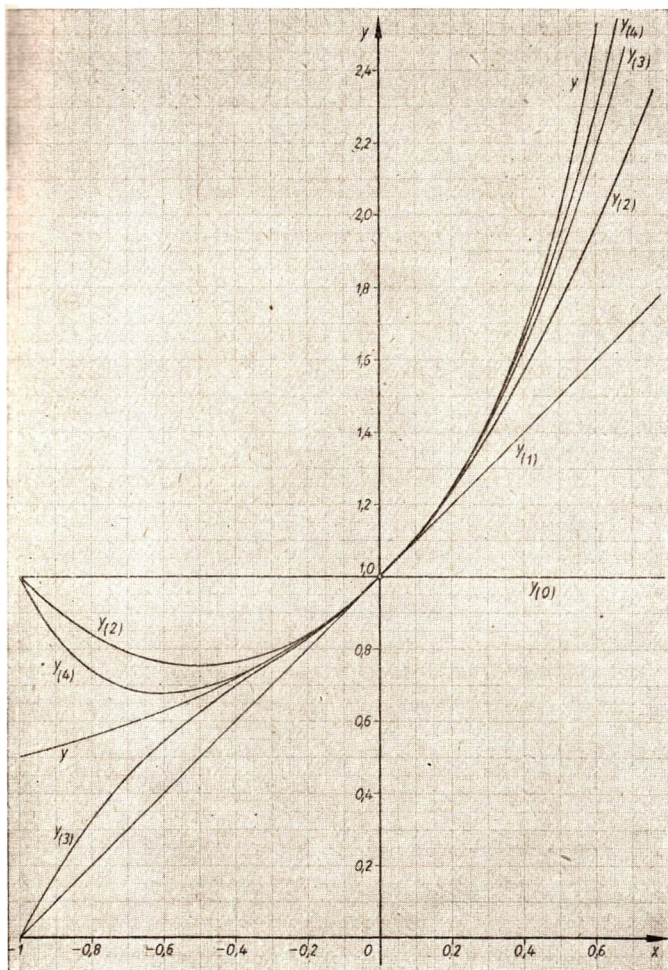


Abb. 10

14. Die logarithmische Reihe

Die Funktion $y = \ln x$ kann an der Stelle $x = 0$ nicht als Potenzreihe dargestellt werden, da die Funktion und ihre Ableitungen an der Stelle $x = 0$ nicht erklärt sind. Wir behandeln daher die Funktion $y = \ln(1+x)$. Ihre Ableitungen sind $y^{(k)} = \frac{(-1)^{k-1} \cdot (k-1)!}{(1+x)^k}$ mit $k = 1, 2, 3, \dots$. Es existieren also alle Ableitungen bis zu beliebiger Ordnung für $x \neq -1$, insbesondere ist

$$y(0) = 0, \quad a^{(k)}(0) = (-1)^{k-1} \cdot (k-1)!.$$

Wir untersuchen das Restglied der Approximation $(n-1)$ -ter Ordnung zur vorgegebenen Funktion. Es ist

$$R_n = \frac{(-1)^{n-1} x^n}{n(1+\vartheta x)^n} \quad \text{mit } 0 < \vartheta < 1.$$

Für alle n und jedes ϑ mit $0 < \vartheta < 1$ ist $\lim_{n \rightarrow \infty} R_n = 0$, wenn nur $|x| < 1$ ist. Für $|x| > 1$ dagegen existiert $\lim_{n \rightarrow \infty} R_n$ nicht.

Die Funktion $y = \ln(1+x)$ ist also unter einschränkenden Bedingungen als die folgende Taylorsche Reihe darstellbar:

$$y = \ln(1+x) = x - \frac{x^2}{2} + \frac{x^3}{3} - \frac{x^4}{4} + \dots = \sum_{k=1}^{\infty} (-1)^{k-1} \frac{x^k}{k}. \quad (6)$$

Sie wird als **logarithmische Reihe** bezeichnet.

Für den Konvergenzradius finden wir

$$r = \lim_{k \rightarrow \infty} \left| \frac{a_k}{a_{k+1}} \right| = \lim_{k \rightarrow \infty} \left| \frac{k+1}{k} \right| = 1,$$

die Reihe ist also für $-1 < x < +1$ absolut konvergent. Für $x = +1$ ergibt sich die nach dem Konvergenzkriterium von Leibniz konvergente Reihe $1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \dots$. Da die Reihe ihrer Absolutbeträge die schon als divergent erkannte harmonische Reihe ist, ist die logarithmische Reihe für $x = +1$ nur bedingt konvergent. Für $x = -1$ erhalten wir die Reihe $-1 - \frac{1}{2} - \frac{1}{3} - \dots = -\left(1 + \frac{1}{2} + \frac{1}{3} + \dots\right)$, die ebenfalls divergent ist. Zusammenfassend stellen wir fest:

Die logarithmische Reihe $\ln(1+x) = \sum_{k=1}^{\infty} (-1)^{k-1} \frac{x^k}{k}$ ist für $|x| < 1$ absolut konvergent, für $x = +1$ bedingt konvergent und für $x = -1$ und $|x| > 1$ divergent.

Die logarithmische Reihe kann zur Berechnung der Logarithmen von Zahlen zwischen 0 und 1 verwendet werden. Aber für diese Aufgabe ist sie recht ungeeignet, wenn $|x|$ wenig von 1 verschieden ist. In diesem Falle konvergiert die Reihe

sehr langsam, es sind also sehr viele Glieder zur Berechnung nötig, wenn eine vorgeschriebene Genauigkeit für den Funktionswert erreicht werden soll. Wir betrachten neben $y = \ln(1+x)$ die Reihe

$$\ln(1-x) = -x - \frac{x^2}{2} - \frac{x^3}{3} - \frac{x^4}{4} - \dots \quad (7)$$

mit dem Konvergenzbereich $-1 \leq x < +1$. Subtrahiert man diese Reihe von der Reihe (6) und beachtet man, daß $\ln a - \ln b = \ln \frac{a}{b}$ ist, so erhält man

$$\ln \frac{1+x}{1-x} = 2 \left(x + \frac{x^3}{3} + \frac{x^5}{5} + \dots \right).$$

Wir setzen $\frac{1+x}{1-x} = z$, also $x = \frac{z-1}{z+1}$ und erhalten damit die Reihe

$$\ln z = 2 \left[\left(\frac{z-1}{z+1} \right) + \frac{1}{8} \left(\frac{z-1}{z+1} \right)^3 + \frac{1}{6} \left(\frac{z-1}{z+1} \right)^5 + \dots \right]. \quad (8)$$

Da die Reihen (6) und (7) für $-1 < x < +1$ absolut konvergent sind, so ist die Reihe (8) für alle positiven $z = \frac{1+x}{1-x}$ absolut konvergent. Es entspricht nämlich wegen der Gleichungen $z(-1) = 0$, $z(+1) = \infty$ und $\frac{dz}{dx} = \frac{2}{(1-x)^2} > 0$ das x -Intervall $-1 < x < +1$ dem z -Intervall $0 < z < \infty$. Die Reihe ermöglicht also für alle positiven Zahlen die Berechnung der Logarithmen. Allerdings ist auch diese Reihe nicht für alle Werte z gleichmäßig gut konvergent. Am brauchbarsten ist die Reihe, wenn z in der Nähe von 1 liegt. Das kann man an der folgenden Tabelle erkennen.

z	0,02	0,5	0,8	1	1,2	1,5	2	11
$\frac{z-1}{z+1}$	$-\frac{49}{51}$	$-\frac{1}{3}$	$-\frac{1}{9}$	0	$+\frac{1}{11}$	$+\frac{1}{5}$	$+\frac{1}{3}$	$+\frac{5}{6}$

Bei der Berechnung der Logarithmen der ganzen Zahlen kann man sich auf die Berechnung der natürlichen Logarithmen der Primzahlen beschränken, denn der Logarithmus einer zusammengesetzten Zahl ist gleich der Summe der Logarithmen ihrer Primfaktoren; so ist zum Beispiel $\ln 15 = \ln 3 + \ln 5$; $\ln 36 = 2 \ln 2 + 2 \ln 3$. Bei größeren Primzahlen konvergiert die Reihe (8) nur langsam. In solchen Fällen versucht man die Zahl, deren Logarithmus zu bestimmen ist, durch eine geeignete Erweiterung so umzuformen, daß zwei oder mehrere leichter zu berechnende Logarithmen entstehen.

Beispiel:

$$a) \ln 97 = \ln \frac{97 \cdot 100}{100} = \ln 100 + \ln \frac{97}{100} = 2 \ln 2 + 2 \ln 5 + \ln \left(1 - \frac{3}{100} \right),$$

b) $\ln 11$ würden nach (8) mit $\frac{z-1}{z+1} = \frac{5}{6}$ nur schlecht konvergieren. Wir setzen darum $\ln 11 = \ln \frac{11 \cdot 9}{9} = \ln 9 + \ln \frac{11}{9} = 2 \ln 3 + \ln \frac{11}{9}$ mit $\frac{z-1}{z+1} = \frac{1}{10}$. Die Berechnung der Zehnerlogarithmen geschieht durch Umrechnung der entsprechenden natürlichen

Logarithmen (vgl. Lehrbuch der Mathematik, 11. Schuljahr, Seite 113). Bekanntlich ist $\lg x = M \ln x$, wobei $M = \frac{1}{\ln 10} = 0,43429 \dots \approx 0,4343$ ist.

15. Die Reihe für $\arctg x$

Wir wenden uns einem letzten Beispiel, der Funktion $y = \arctg x$ zu und beziehen uns nur auf die Hauptwerte dieser Funktion. Ihre erste Ableitung ist $y' = \frac{1}{1+x^2}$. Die Bestimmung der weiteren Ableitungen wird mit wachsender Ordnung der Ableitungen immer unübersichtlicher. Wir werden daher einen anderen Weg zur Entwicklung der Funktion $y = \arctg x$ an der Stelle $x=0$ einschlagen. Die abgeleitete Funktion $y' = \frac{1}{1+x^2}$ läßt sich mit Hilfe der binomischen Reihe als Potenzreihe schreiben. Es ist

$$y' = \frac{1}{1+x^2} = (1+x^2)^{-1} = 1 - x^2 + x^4 - x^6 + - \dots$$

mit dem Konvergenzintervall $-1 < x < +1$. In diesem Intervall ist die angegebene Reihe absolut konvergent, und als solche kann sie gliedweise integriert werden, sofern nur die Integrationsgrenzen dem Konvergenzbereich der zu integrierenden Reihe angehören. Die untere Grenze sei 0, die obere x . Dann ergibt sich aus der Reihe $1 - x^2 + x^4 - x^6 + - \dots$ durch Integration die Reihe

$$\int_0^x (1 + \xi^2)^{-1} d\xi = \int_0^x (1 - \xi^2 + \xi^4 - + \dots) d\xi,$$

mithin

$$y = \arctg x = x - \frac{x^3}{3} + \frac{x^5}{5} - \frac{x^7}{7} + - \dots \quad (9)$$

Ihr Konvergenzradius ist $r=1$, wie man leicht nachweisen kann. Die Reihe ist auch noch für $x = \pm 1$ bedingt konvergent, denn in diesem Falle erhalten wir die alternierende Reihe $\pm \left(1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} + - \dots\right)$, bei der die absoluten Beträge der Glieder eine monotone Nullfolge bilden und die deshalb nach dem Leibnizschen Konvergenzkriterium für alternierende Reihen konvergiert. Ob allerdings die für $x = \pm 1$ erhaltenen bedingt konvergenten Reihen tatsächlich die Funktionswerte $\arctg(\pm 1) = \pm \frac{\pi}{4}$ darstellen, geht aus unserer Ableitung der Reihe (9) nicht hervor, da die Integration nur unter der Voraussetzung $-1 < x < +1$ möglich ist. Die Übereinstimmung besteht jedoch für jedes x , was wir aber nicht beweisen wollen.

Die Reihe

$$\frac{\pi}{4} = 1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} + \frac{1}{9} - \frac{1}{11} + - \dots \quad (10)$$

wird auch Leibnizsche Reihe genannt. Zur numerischen Berechnung von π ist sie wegen ihrer schlechten Konvergenz wenig geeignet. Zur Bestimmung von nur zwei Dezimalen wären 300 Glieder notwendig.

Mit $x = \frac{1}{3} \sqrt{3}$ ergibt sich $\operatorname{arctg} \left(\frac{1}{3} \sqrt{3} \right) = \frac{\pi}{6}$. Damit erhalten wir die Reihe

$$\frac{\pi}{6} = \frac{1}{3} \sqrt{3} \cdot \left(1 - \frac{1}{3} \cdot \frac{1}{3} + \frac{1}{5} \cdot \frac{1}{3^3} - \frac{1}{7} \cdot \frac{1}{3^5} + \frac{1}{9} \cdot \frac{1}{3^7} - + \dots \right). \quad (11)$$

Setzt man $\frac{\pi}{4} = \operatorname{arctg} \frac{1}{p} + \operatorname{arctg} \frac{1}{q}$, $\operatorname{arctg} \frac{1}{p} = \alpha$, $\operatorname{arctg} \frac{1}{q} = \beta$, so ist $\alpha + \beta = \frac{\pi}{4}$. Folglich ist $\operatorname{tg} (\alpha + \beta) = \operatorname{tg} \frac{\pi}{4} = 1$. Nun ist nach dem Additionstheorem der Tangensfunktion $\operatorname{tg} (\alpha + \beta) = \frac{\operatorname{tg} \alpha + \operatorname{tg} \beta}{1 - \operatorname{tg} \alpha \operatorname{tg} \beta}$. Also ist

$$\frac{\frac{1}{p} + \frac{1}{q}}{1 - \frac{1}{p} \cdot \frac{1}{q}} = 1 \quad \text{oder} \quad \frac{q + p}{p \cdot q - 1} = 1.$$

Wir lösen diese Gleichung nach q auf und erhalten $q = \frac{p+1}{p-1}$. Setzen wir $p = 2$, so ergibt sich $q = 3$, das heißt, es ist $\frac{\pi}{4} = \operatorname{arctg} \frac{1}{2} + \operatorname{arctg} \frac{1}{3}$. Wir erhalten also zur Berechnung von $\frac{\pi}{4}$ die Doppelreihe

$$\frac{\pi}{4} = \left(\frac{1}{2} - \frac{1}{3 \cdot 2^3} + \frac{1}{5 \cdot 2^5} - \dots \right) + \left(\frac{1}{3} - \frac{1}{3 \cdot 3^3} + \frac{1}{5 \cdot 3^5} - \dots \right). \quad (12)$$

Eine weitere Zerlegung dieser Art ist $\frac{\pi}{4} = \operatorname{arctg} \frac{1}{2} + \operatorname{arctg} \frac{1}{5} + \operatorname{arctg} \frac{1}{8}$. Man erhält sie, indem man $\operatorname{arctg} \frac{1}{3}$ zerlegt in $\operatorname{arctg} \frac{1}{r} + \operatorname{arctg} \frac{1}{s}$. Daraus ergibt sich $s = \frac{3r+1}{r-3}$, mit $r = 5$ und $s = 8$ erhält man dann die Zerlegung für $\frac{\pi}{4}$.

Aufgaben

1. Der Konvergenzradius für die Potenzreihe $G(x) = 1 + x + x^2 + x^3 + \dots = \sum_{\nu=0}^{\infty} x^{\nu}$ ist zu bestimmen. Welche Funktion stellt die Reihe innerhalb ihres Konvergenzintervalls dar?

2. Es ist zu untersuchen, welche Funktion die Reihe

$$1 + 2x + 3x^2 + 4x^3 + \dots = \sum_{\nu=0}^{\infty} (\nu+1) \cdot x^{\nu}$$

3. Es ist die Reihe $x + \frac{1}{2}x^2 + \frac{1}{3}x^3 + \frac{1}{4}x^4 + \dots = \sum_{\nu=1}^{\infty} \frac{x^{\nu}}{\nu}$ zu differenzieren und das Ergebnis mit der Reihe in Aufgabe 1 zu vergleichen.

4. Es ist der Konvergenzradius der Potenzreihen

$$\text{a) } \sum_{\nu=0}^{\infty} \frac{1}{\nu!} x^{\nu} = 1 + x + \frac{x^2}{2!} + \frac{x^3}{3!} + \dots \quad \text{b) } \sum_{\nu=0}^{\infty} \nu! x^{\nu} = 1 + x + 2!x^2 + 3!x^3 + \dots$$

zu bestimmen.

5. Begründen Sie, weshalb sich die Funktionen

$$\text{a) } f(x) = \ln x, \quad \text{b) } f(x) = \frac{1}{x^2}, \quad \text{c) } f(x) = \frac{1}{\sqrt{x}}, \quad \text{d) } f(x) = \frac{1}{\sin x}$$

an der Stelle $x = 0$ nicht in eine Taylorreihe entwickeln lassen!

6. Es sind für die Potenzreihen

$$\text{a) } \sum_{k=0}^{\infty} (-1)^k x^k, \quad \text{b) } \sum_{k=1}^{\infty} (-1)^{k-1} \frac{x^k}{k}$$

die Näherungspolynome bis zur 4. Ordnung zu berechnen und die zugehörigen Schmiegungsparabeln zu zeichnen.

7. Erläutern Sie, weshalb die zu den Kurven in Abbildung 11 gehörenden Taylorschen Reihen an der Stelle $x_0 = 0$ nicht mehr konvergent sein können!

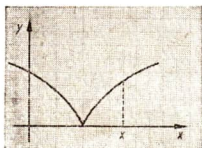


Abb. 11 a

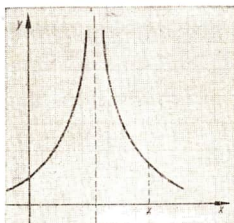


Abb. 11 b

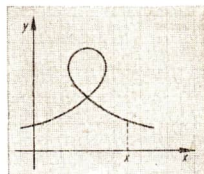


Abb. 11 c

8. Welche Voraussetzungen fehlen bei den Funktionen

$$\text{a) } f(x) = 3 + \sqrt{x}, \quad \text{b) } f(x) = \sqrt[3]{x^2}, \quad \text{c) } f(x) = x\sqrt{x}$$

für die Entwicklung einer Taylorschen Reihe an der Stelle $x = 0$?

9. Berechnen Sie die Binomialkoeffizienten

$$\begin{aligned} \text{a) } & \binom{5}{3}, \binom{8}{8}, \binom{7}{9}, \binom{-1}{3}, \binom{-3}{5}, \binom{-6}{6}, \binom{-7}{4}; \\ \text{b) } & \binom{1}{\frac{1}{2}}, \binom{5}{\frac{5}{2}}, \binom{4}{\frac{4}{3}}, \binom{1}{\frac{1}{4}}, \binom{3}{\frac{3}{4}}, \binom{6}{\frac{6}{3}}, \binom{3}{\frac{3}{8}}; \\ \text{c) } & \binom{-1}{\frac{1}{2}}, \binom{-1}{\frac{1}{4}}, \binom{-2}{\frac{2}{3}}, \binom{-3}{\frac{3}{2}}, \binom{-3}{\frac{3}{4}}, \binom{-1}{\frac{1}{3}}, \binom{-3}{\frac{3}{2}}. \end{aligned}$$

10. Die allgemeine binomische Formel ist für

$$\text{a) } (a + b)^n, \quad \text{b) } (a - b)^n$$

anzugeben.

11. Lesen Sie aus dem binomischen Lehrsatz Summenformeln für die Binomialkoeffizienten ab, indem Sie für x die Werte ± 1 einsetzen!

12. Es sind die Funktionen

$$\begin{array}{lll} \text{a)} y = (1+x)^{-3}, & \text{b)} y = (1-x)^{-5}, & \text{c)} y = \frac{1}{(1-x)^4}, \\ \text{d)} y = (1+2x)^{-2}, & \text{e)} y = (1-x^2)^{-6}, & \text{f)} y = \left(1 + \frac{1}{2}x\right)^{-4}, \\ \text{g)} y = \sqrt[3]{(1-x)^3}, & \text{h)} y = \sqrt[5]{(1+x)^2}, & \text{i)} y = (1-3x)^{\frac{1}{2}}, \\ \text{k)} y = \left(1 - \frac{1}{2}x\right)^{\frac{1}{3}}, & \text{l)} y = (4+x)^{\frac{5}{2}}, & \text{m)} y = \sqrt{a+bx}, \\ \text{n)} y = (1+x^2)^{-\frac{1}{2}}, & \text{o)} y = \sqrt{a^2-x^2}, & \text{p)} y = \frac{1}{\sqrt[3]{\left(1 - \frac{1}{3}x\right)^2}} \end{array}$$

an der Stelle $x_0 = 0$ zu entwickeln. Zu jeder Taylorschen Reihe sind das allgemeine Glied und der Konvergenzradius anzugeben.

13. Es sind die Funktionen

$$\text{a)} y = e^{1+x}, \quad \text{b)} y = e^{3x}, \quad \text{c)} y = e^{x^2}, \quad \text{d)} y = e^{-x^2}$$

als Taylorsche Reihen darzustellen.

14. Die Funktionen

$$\text{a)} y = \frac{1}{2}(e^x + e^{-x}) \quad (-\operatorname{Cof} x), \quad \text{b)} y = \frac{1}{2}(e^x - e^{-x}) \quad (-\operatorname{Sin} x)$$

sind in Reihen zu entwickeln (vgl. Lehrbuch der Mathematik, 11. Schuljahr, Seite 125). Es ist nachzuprüfen, ob die beiden Funktionen gerade oder ungerade sind und welche Symmetrieverhältnisse die zugehörigen Kurven aufweisen. Die Reihen sind mit der Sinusreihe bzw. Kosinusreihe zu vergleichen.

15. Entwickeln Sie die Funktion

$$\text{a)} y = \sin 2x, \quad \text{b)} y = \sin \frac{1}{3}x, \quad \text{c)} y = \sin ax$$

Vergleichen Sie die erhaltenen Reihen mit der Reihe für $y = \sin x$ und leiten Sie aus dem Vergleich eine Regel ab!

16. Mit Hilfe der in Aufgabe 15 abgeleiteten Regel sind die Reihen für

$$\text{a)} y = \sin \frac{2}{3}x, \quad \text{b)} y = \sin 5x, \quad \text{c)} y = \sin(x^2)$$

anzugeben.

17. Es sind die Funktionen

$$\text{a)} y = \cos 3x, \quad \text{b)} y = \cos \frac{1}{2}x, \quad \text{c)} y = \cos(ax)$$

in eine Reihe zu entwickeln. Die so erhaltenen Reihen sind mit der Kosinusreihe zu vergleichen. Aus diesem Vergleich ist eine Regel abzuleiten.

18. Geben Sie mit Hilfe der in Aufgabe 17 abgeleiteten Regel die Reihe für

a) $y = \cos \frac{1}{3} x$, b) $y = \cos \frac{4}{3} x$, c) $y = \cos(x^2)$, d) $y = \cos(\sqrt{x})$
an!

19. Es sind die Funktionen $y = \sin^2 x$ und $y = \frac{1}{2}(1 - \cos 2x)$ in eine Taylorreihe zu entwickeln. Auf welche bekannte trigonometrische Formel führt der Vergleich beider Reihen?

20. Die Funktionen $y = \sin 2x$ und $y = 2 \sin x \cdot \cos x$ sind in eine Taylorreihe zu entwickeln und miteinander zu vergleichen. Welche Formel läßt sich daraus ableiten?

21. Bestimmen Sie durch Multiplikation der Sinus- und Kosinusreihe die Reihe für $y = 2 \sin x \cdot \cos x$! Warum kann man die beiden Reihen wie endliche Polynome multiplizieren? Vergleichen Sie die Reihe mit der von Aufgabe 15a!

Berechnen Sie mit Hilfe der Sinus- und Kosinusreihe die Funktionswerte für y auf vier Dezimalen!

22. a) $\sin 1$, b) $\sin 0,1$, c) $\sin 0,04$, d) $\sin \frac{\pi}{3}$, e) $\sin \frac{\pi}{8}$!

23. a) $\cos 0,01$, b) $\cos 1$, c) $\cos 1,5$, d) $\cos \frac{\pi}{3}$, e) $\cos \frac{\pi}{2}$.

Welche Winkel gehören zu diesen Entwicklungen der Aufgaben 22 und 23?

24. a) $\sin 1^\circ$, b) $\sin 4^\circ$, c) $\sin 2^\circ 30'$, d) $\sin 30^\circ$, e) $\sin 90^\circ$.

25. a) $\cos 15'$, b) $\cos 2^\circ$, c) $\cos 10^\circ$, d) $\cos 45^\circ$, e) $\cos 90^\circ$.

Anmerkung: In den Aufgaben 24 und 25 ist das Gradmaß in Bogenmaß umzurechnen.

26. a) $\sin 10^\circ$, b) $\sin 25^\circ$, c) $\cos 15^\circ$, d) $\cos 20^\circ$.

27. a) $\sin 200^\circ$, b) $\sin 425^\circ$, c) $\sin(-800^\circ)$.

28. a) $\cos 5^\circ$, b) $\cos 7,5^\circ$, c) $\cos 12,5^\circ$.

29. a) $\cos 300^\circ$, b) $\cos 555^\circ$, c) $\cos(-700^\circ)$!

Anmerkung: In den Aufgaben 27 und 29 sind die Winkel auf den ersten Quadranten zu reduzieren.

30. Es sind die natürlichen Logarithmen der Primzahlen von 2 bis 7 zu berechnen.

31. Die folgenden Werte sind mit Hilfe der Reihe für $\ln z$ zu berechnen:

a) $\ln 2$; b) $\ln 1,2$; c) $\ln 1,5$; d) $\ln 5$; e) $\ln 11$; f) $\ln 23$; g) $\ln 51$.

32. Bestimmen Sie die natürlichen Logarithmen von

a) 72, b) 1000, c) 1327, d) 1323, e) 7,5, f) 30,87,

indem Sie die Logarithmen in eine Summe logarithmierter Primzahlen zerlegen!

33. Die natürlichen Logarithmen der Primzahlen

a) 19; b) 37; c) 61; d) 73; e) 97; f) 211; g) 197

sind durch geeignete Umformung der Logarithmanden zu bilden.

34. Berechnen Sie die Zehnerlogarithmen der Logarithmanden von Aufgabe 30!

35. In der Zinseszins- und Rentenrechnung werden häufig große Vielfache der Logarithmen zu den Aufzinsungsfaktoren $r = 1 + \frac{p}{100}$ gebraucht. Für diese Rechnungen führt die Verwendung der üblichen vierstelligen Logarithmentafeln zu erheblichen Ungenauigkeiten. Deshalb verwendet man für diese Rechnungen Logarithmentafeln mit höherer Stellenzahl.

Der Prozentsatz p möge zwischen 1% und 5% liegen. Es ist mit Hilfe einer Restgliedabschätzung zu bestimmen, bis zu welcher Ordnung die Funktion $\ln r = \ln(1 + x)$ approximiert werden muß, damit eine Genauigkeit von 6 Dezimalen für $\ln r$ gewährleistet ist.

36. Es ist die ungleich gute Konvergenz der $\operatorname{arc} \operatorname{tg}$ -Reihe zu untersuchen, indem für $x = 1; 0,1; 0,01$ das Glied bestimmt wird, das bei einer vorgegebenen Genauigkeit von $\varepsilon = 10^{-4}$ keinen Beitrag mehr zur vierten Dezimale der Summe gibt.

37. Es ist π mit Hilfe der Reihe für

a) $\frac{\pi}{6}$ (Reihe 11), b) $\frac{\pi}{4}$ (Reihe 10)

auf 5 Dezimalen genau zu berechnen.

38. Es ist die Reihe für die Funktion $y = \arcsin x$ aus der Reihe ihrer Ableitung $y' = \frac{1}{\sqrt{1-x^2}}$ zu bestimmen.

39. Leiten Sie aus der $\arcsin$ -Reihe eine Reihe für $\frac{\pi}{6}$ ab!

V. Zusammenstellung der wichtigsten Sätze und Reihenentwicklungen

Der Mittelwertsatz der Differentialrechnung:

$$\frac{f(b) - f(a)}{b - a} = f'(\xi) \quad \text{mit } a \leq \xi \leq b.$$

Ist speziell $f(b) = f(a)$, so ergibt sich der Satz von Rolle:

$$0 = f'(\xi) \quad \text{mit } 0 \leq \xi \leq b.$$

Der verallgemeinerte Mittelwertsatz der Differentialrechnung:

$$\frac{f(x+h) - f(x)}{g(x+h) - g(x)} = \frac{f'(x+\vartheta h)}{g'(x+\vartheta h)} \quad \text{mit } 0 < \vartheta < 1.$$

Die Taylorsche Formel:

$$G(x_0 + h) = G(x_0) + \frac{h}{1!} G'(x_0) + \frac{h^2}{2!} G''(x_0) + \cdots + \frac{h^n}{n!} G^{(n)}(x_0).$$

Die Maclaurinsche Formel:

$$G(x) = G(0) + \frac{x}{1!} G'(0) + \frac{x^2}{2!} G''(0) + \cdots + \frac{x^n}{n!} G^{(n)}(0).$$

Der Taylorsche Lehrsatz:

$$f(x_0 + h) = f(x_0) + \frac{h}{1!} f'(x_0) + \frac{h^2}{2!} f''(x_0) + \cdots + \frac{h^{n-1}}{(n-1)!} f^{(n-1)}(x_0) + R_n$$

mit $R_n = \frac{h^n}{n!} f^{(n)}(x_0 + \vartheta h)$, wobei $0 < \vartheta < 1$ (Restglied von Lagrange).

Die Taylorsche Reihe:

$$\text{Es ist } f(x_0 + h) = f(x_0) + \frac{h}{1!} f'(x_0) + \frac{h^2}{2!} f''(x_0) + \cdots,$$

wenn $\lim_{n \rightarrow \infty} R_n(x_0; h) = 0$ gilt.

Die Exponentialreihe:

$$e^x = 1 + \frac{1}{1!} x + \frac{1}{2!} x^2 + \frac{1}{3!} x^3 + \cdots.$$

Die Reihen für $\sin x$ und $\cos x$:

$$\sin x = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \cdots,$$

$$\cos x = 1 - \frac{x^2}{2!} + \frac{x^4}{4!} - \frac{x^6}{6!} + \cdots.$$

Die binomische Reihe:

$$(a+x)^\beta = a^\beta + \binom{\beta}{1} a^{\beta-1} x + \binom{\beta}{2} a^{\beta-2} x^2 + \binom{\beta}{3} a^{\beta-3} x^3 + \cdots.$$

Die logarithmische Reihe:

$$\ln(1+x) = x - \frac{x^2}{2} + \frac{x^3}{3} - \frac{x^4}{4} + \cdots = \sum_{k=1}^{\infty} (-1)^{k-1} \frac{x^k}{k},$$

$$\ln z = 2 \left[\left(\frac{z-1}{z+1} \right) + \frac{1}{3} \left(\frac{z-1}{z+1} \right)^3 + \frac{1}{5} \left(\frac{z-1}{z+1} \right)^5 + \cdots \right] \quad \text{mit } z = \frac{1+x}{1-x}.$$

C. Sphärische Trigonometrie

Bei den bisherigen Betrachtungen in der Trigonometrie sind wir davon ausgegangen, daß die Figuren in einer Ebene liegen. Wir wollen in den folgenden Abschnitten die trigonometrischen Untersuchungen auf die Kugeloberfläche ausdehnen. Dieses Teilgebiet der Mathematik, mit dem wir uns jetzt beschäftigen werden, heißt **sphärische Trigonometrie**. Die Bezeichnung ist von dem griechischen Wort *sphaira* (Kugel) abgeleitet.

I. Grundbegriffe der Kugelgeometrie

1. Großkreise und Kleinkreise, kürzeste Entfernung zweier Punkte

Wird eine Kugel von einer Ebene geschnitten, so entsteht als Schnittfigur ein Kreis. Die Größe des Schnittkreises ist vom Abstände a der Ebene vom Mittelpunkt der Kugel abhängig. Es ergeben sich die folgenden Möglichkeiten:

1. $a = 0$ Der Mittelpunkt des Schnittkreises fällt mit dem Mittelpunkt der Kugel zusammen. Der Radius des Schnittkreises ist gleich dem Radius der Kugel. Es ist offensichtlich, daß es keinen Schnittkreis mit größerem Radius geben kann. Man bezeichnet deshalb diesen Schnittkreis als **Großkreis**.
2. $0 < a < r$ Der Radius des Schnittkreises ist $\rho = \sqrt{r^2 - a^2}$. (Dies kann mit Hilfe einer Zeichnung leicht hergeleitet werden.) Es ist also in diesem Falle ρ stets kleiner als r . Man bezeichnet deshalb diese Schnittkreise als **Kleinkreise**.
3. $a = r$ Die Ebene berührt die Kugel. Der Schnittkreis entartet zu einem Punkt.
4. $a > r$ Die Ebene schneidet die Kugel nicht. Es entsteht also kein Schnittkreis.

In der Ebene ist die Strecke die kürzeste Verbindung zweier Punkte. Da es auf der Kugeloberfläche keine Geraden gibt, tritt hier an die Stelle der Strecke der kleinere Bogen eines Großkreises durch die beiden Punkte. Dieser Bogen ist auf der Kugeloberfläche die kürzeste Verbindung zwischen zwei Punkten.

Es soll hier nur bewiesen werden, daß der kleinere Bogen des Großkreises unter allen möglichen Kreisbogen auf der Kugeloberfläche die kürzeste Verbindung zweier Punkte ist. Um zu beweisen, daß er die kürzeste Verbindung unter allen überhaupt möglichen Kurven ist, muß man Hilfsmittel heranziehen, die über den Lehrstoff der Oberschule hinausgehen.

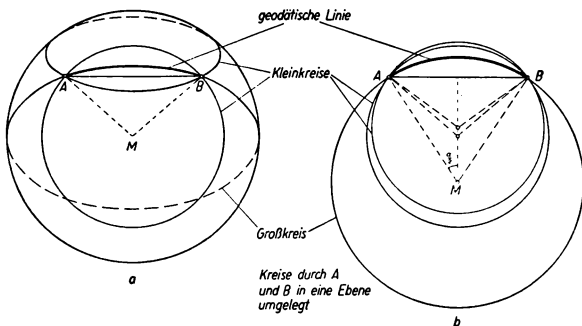


Abb. 12

In der Abbildung 12a sind drei Kreise durch die Punkte A und B der Kugeloberfläche gelegt. Die Abbildung 12b zeigt die Umlegung dieser Kreise in eine Ebene. Die drei Kreise haben die Sehne $AB = s$ gemeinsam. Die Länge b des Bogens wird durch die Formel $b = \hat{a} \cdot r$ bestimmt. Der Winkel α kann aus $\sin \frac{\alpha}{2} = \frac{s}{2r}$ zu $\alpha = 2 \arcsin \frac{s}{2r}$ berechnet werden. Dann gilt also $b = 2r \arcsin \frac{s}{2r}$.

Für $\arcsin x$ gilt die Reihendarstellung

$$\arcsin x = x + \frac{1}{2} \cdot \frac{x^3}{3} + \frac{1 \cdot 3}{2 \cdot 4} \cdot \frac{x^5}{5} + \dots \quad \text{mit } |x| \leq 1,$$

es ist also

$$= \frac{s}{2r} + \frac{1}{2} \cdot \frac{s^3}{3(2r)^3} + \frac{1 \cdot 3}{2 \cdot 4} \cdot \frac{s^5}{5(2r)^5} + \dots$$

wobei $\left| \frac{s}{2r} \right| \leq 1$ vorausgesetzt werden muß. Diese Voraussetzung ist tatsächlich erfüllt,

es gilt sogar $0 < \frac{s}{2r} \leq 1$, denn s und r sind stets positiv, und $2r$ kann nicht kleiner als s werden.

Man erhält damit

$$b = 2r \left[\frac{s}{2r} + \frac{1}{2} \cdot \frac{s^3}{3(2r)^3} + \frac{1 \cdot 3}{2 \cdot 4} \cdot \frac{s^5}{5(2r)^5} + \dots \right]$$

oder

$$b = s + \frac{1}{2} \cdot \frac{s^3}{3(2r)^2} + \frac{1 \cdot 3}{2 \cdot 4} \cdot \frac{s^5}{5(2r)^4} + \dots$$

Da s konstant ist, werden mit wachsendem r vom zweiten Glied an alle Glieder dieser Reihe und damit auch ihre Summe immer kleiner. Das bedeutet aber, daß der Bogen b um so kleiner wird, je größer der Radius r wird.

Nach Definition ist auf der Kugeloberfläche der Großkreis unter allen möglichen Kreisen derjenige mit dem größten Radius, also ist der kleinere Bogen des Großkreises unter allen möglichen Kreisbogen die kürzeste Verbindung zwischen A und B .

Den kleineren Bogen des Großkreises durch zwei Punkte auf der Kugeloberfläche bezeichnen wir als **geodätische Linie**. (Allgemein wird die kürzeste Verbindung zweier Punkte auf jeder gekrümmten Fläche als geodätische Linie bezeichnet.)

Die Länge der geodätischen Linie kann man auf zwei Arten messen: 1. mit einem Längenmaß auf dem kleineren Bogen des Großkreises; 2. als Winkel, den die beiden Radien zu den Endpunkten der geodätischen Linie bilden. Im ersten Falle ist die Länge der geodätischen Linie außer von der Lage der Punkte auch vom Kugelradius abhängig, im zweiten Falle dagegen nicht. Zur besseren Unterscheidung bezeichnen wir die Länge der geodätischen Linie als **sphärischen Abstand**, wenn sie im Längenmaß gemessen ist, dagegen als **sphärische Länge**, wenn sie im Winkelmaß gemessen ist. Für die Länge der geodätischen Linie gilt die folgende Relation:

$$b = r \cdot \hat{a} = r \cdot \frac{\pi \cdot \alpha}{180^\circ}.$$

Dabei ist α die sphärische Länge und b der sphärische Abstand (Abb. 13).

Für die Erde ist zum Beispiel die Länge der geodätischen Linie, die zur sphärischen Länge 1° gehört, annähernd 111 km.

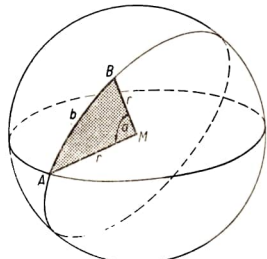


Abb. 13

2. Das sphärische Zweieck, Winkel am Zweieck, Flächeninhalt

Im allgemeinen wird durch zwei Punkte auf der Kugel ein Großkreis eindeutig bestimmt, da der Mittelpunkt der Kugel mit diesen beiden Punkten zusammen die Schnittebene festlegt. Sind speziell die beiden Punkte Endpunkte eines Durchmessers, so liegen sie zusammen mit dem Mittelpunkt auf einer Geraden. In diesem Fall lassen sich durch diese beiden Punkte, die diametral zueinander liegen, beliebig viele Großkreise ziehen, die sämtlich einander halbieren.

Durch zwei nicht zusammenfallende Großkreise wird die Kugeloberfläche in vier Teile zerlegt, von denen je zwei einander kongruent sind (Abb. 14). Die so entstandenen Teile der Kugelfläche nennt man **Kugelzweiecke**. Alle Kugelzweiecke sind gleichseitig. Die sphärische Länge ihrer Seiten ist 180° . Die Größe des Kugelzweiecks ist von der Größe des Winkels α

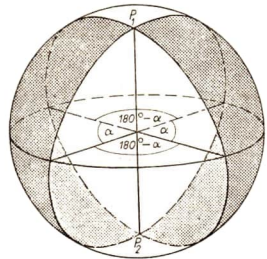


Abb. 14

zwischen den erzeugenden Großkreisebenen abhängig. Diesen Winkel α bezeichnen wir als den Winkel des Kugelzweiecks.

Der Flächeninhalt eines Kugelzweiecks läßt sich aus der Proportion

$$F : O = F : 4 \pi r^2 = \alpha : 360^\circ$$

als

$$F = \frac{4 \pi r^2 \alpha}{360^\circ} = \frac{\pi r^2 \alpha}{90^\circ}$$

bestimmen, wobei F die Fläche des Zweiecks und O die Oberfläche der Kugel bedeuten.

Wird der Winkel jedoch im Bogenmaß gemessen, so ist

$$F : O = F : 4 \pi r^2 = \hat{\alpha} : 2 \pi,$$

also

$$F = \frac{4 \pi r^2 \hat{\alpha}}{2 \pi} = 2 r^2 \hat{\alpha}.$$

3. Das sphärische Dreieck, Flächeninhalt, Winkelsumme

Wir wollen im folgenden Dreiecke auf der Kugeloberfläche betrachten. Ein Kugeldreieck ist die Figur, die von drei Großkreisbogen begrenzt wird. Die Großkreisbogen sind dann nach Abschnitt 1 die Seiten des Dreiecks. Die Winkel zwischen den erzeugenden Großkreisebenen sind entsprechend Abschnitt 2 die Winkel des Dreiecks. Wir beschränken uns bei den folgenden Untersuchungen auf solche Dreiecke, bei denen jede Seite kleiner als 180° ist. Diese Dreiecke werden als **Eulersche Dreiecke** bezeichnet.

Die vollständigen Großkreise schneiden sich paarweise noch einmal in den drei Punkten A_1 , B_1 und C_1 , die jeweils diametral den Punkten A , B und C gegenüberliegen (Abb. 15). Dadurch entsteht ein flächengleiches, symmetrisches sphärisches Dreieck $A_1B_1C_1$, das aber den entgegengesetzten Umlaufsinn des ursprünglichen Dreiecks ABC hat. Dieses Dreieck heißt das **Scheiteldreieck** zum Dreieck ABC . Neben dem Dreieck ABC und seinem Scheiteldreieck $A_1B_1C_1$ entstehen aber auf der Kugel noch weitere sechs Dreiecke:

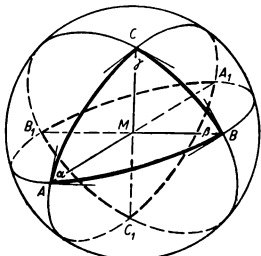


Abb. 15

AB_1C_1 , A_1BC , A_1B_1C , AB_1C , A_1B_1C und ABC_1 .

Die Flächen der acht Dreiecke bedecken vollständig die Kugeloberfläche, ohne einander zu überschneiden.

Da je zwei gegenüberliegende Dreiecke flächengleich sind, genügt es, wenn wir uns bei den folgenden Betrachtungen auf die Halbkugel beschränken. Wir bezeichnen den Flächeninhalt des Dreiecks ABC mit F , den des Dreiecks A_1BC mit F_1 , den des Dreiecks AB_1C mit F_2 und den des Dreiecks A_1B_1C mit F_3 . Es gilt also

$$F + F_1 + F_2 + F_3 = 2\pi r^2. \quad (1)$$

Andererseits bilden die Flächen F und F_1 ein Kugelzweieck mit dem Winkel α , die Flächen F und F_2 ein solches mit dem Winkel β und die Fläche des Scheiteldreiecks A_1B_1C , die der Größe nach gleich F ist, und F_3 ein Kugelzweieck mit dem Winkel γ . Daher ist nach Abschnitt 2

$$F + F_1 = \frac{\pi r^2 \alpha}{90^\circ},$$

$$F + F_2 = \frac{\pi r^2 \beta}{90^\circ},$$

$$F + F_3 = \frac{\pi r^2 \gamma}{90^\circ}.$$

Durch Addition dieser drei Gleichungen ergibt sich

$$3F + F_1 + F_2 + F_3 = \frac{\pi r^2}{90^\circ} (\alpha + \beta + \gamma)$$

und nach Subtraktion von (1)

$$2F = \frac{\pi r^2}{90^\circ} (\alpha + \beta + \gamma) - \frac{180^\circ \pi r^2}{90^\circ},$$

$$2F = \frac{\pi r^2}{90^\circ} (\alpha + \beta + \gamma - 180^\circ)$$

oder

$$F = \frac{\pi r^2}{180^\circ} (\alpha + \beta + \gamma - 180^\circ). \quad (2)$$

Wird $\alpha + \beta + \gamma - 180^\circ = \varepsilon$ gesetzt, dann wird

$$F = \frac{\pi r^2}{180^\circ} \varepsilon. \quad (3)$$

Die Größe ε heißt **sphärischer Exzeß**.

Der Flächeninhalt spielt bei praktischen Berechnungen kaum eine Rolle. Die Bedeutung der Formeln (2) und (3) liegt nicht so sehr in der Tatsache, daß sie die Berechnung des Flächeninhaltes eines sphärischen Dreiecks ermöglichen, sondern vielmehr darin, daß man aus ihnen eine Aussage über die Winkelsumme im sphärischen Dreieck herleiten kann.

Voraussetzung für die Entstehung eines Dreiecks ist, daß die Punkte A , B und C nicht auf demselben Großkreise liegen. Der Flächeninhalt des durch diese Punkte bestimmten Dreiecks ist stets größer als Null. Dann ist nach (2) und (3)

$$\varepsilon = \alpha + \beta + \gamma - 180^\circ > 0,$$

also

$$\alpha + \beta + \gamma > 180^\circ.$$

Die Winkelsumme im sphärischen Dreieck ist demnach stets größer als 180° . Die Größe ε gibt dabei den Überschuß über 180° an. Dadurch wird die Bezeichnung sphärischer Exzeß verständlich.

Andererseits ist die Fläche eines Eulerschen Dreiecks stets kleiner als die der Halbkugel. Es besteht mithin auch die Beziehung

$$\frac{\pi r^2}{180^\circ}(\alpha + \beta + \gamma - 180^\circ) < 2\pi r^2,$$

aus der unmittelbar folgt

$$\alpha + \beta + \gamma - 180^\circ < 360^\circ,$$

$$\alpha + \beta + \gamma < 540^\circ.$$

Für die Winkelsumme im sphärischen Dreieck ergibt sich damit die folgende Ungleichung:

$$180^\circ < \alpha + \beta + \gamma < 540^\circ$$

oder

$$\pi < \hat{\alpha} + \hat{\beta} + \hat{\gamma} < 3\pi.$$

Im sphärischen Dreieck ist die Winkelsumme stets größer als 180° und kleiner als 540° . Der sphärische Exzeß ist stets größer als 0° und kleiner als 360° .

4. Das sphärische Dreieck und die körperliche Ecke

Stoßen n Kanten ($n > 2$), von denen keine drei in einer Ebene liegen, in einem Punkt zusammen, so bilden sie eine n -seitige **körperliche Ecke**.

Durch die drei Seiten eines Eulerschen Dreiecks und den Mittelpunkt der Kugel werden drei Ebenen bestimmt, die sich paarweise in drei Geraden schneiden. Diese drei Geraden sind die Kanten einer dreiseitigen körperlichen Ecke. Es besteht also ein Zusammenhang zwischen dem sphärischen Dreieck und der dreiseitigen körperlichen Ecke.

Die Abbildung 16a zeigt das räumliche Bild einer dreiseitigen körperlichen Ecke, die zu einem sphärischen Dreieck ABC ergänzt ist. Die Projektion des Punktes A auf die Ebene MBC sei A' . Wir fällen von A auf die Kanten MB und MC die Lote. Ihre Fußpunkte seien D und E . Nun verbinden wir D und E mit A' . Dann ist der Winkel ADA' , der Neigungswinkel β der Ebene MBA gegen die Ebene MBC , laut Definition der Winkel β des sphärischen Dreiecks. Entsprechend ist der Winkel AEA' der Winkel γ .

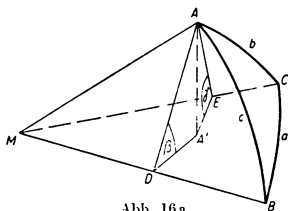


Abb. 16a

Wir wollen nun die wahren Größen der Seiten a , b und c und der Winkel β und γ konstruieren. Dazu wählen wir eine der drei Seitenflächen der körperlichen Ecke, zum Beispiel MBC , als Grundrißebene; die beiden anderen Seitenflächen, also MAB und MAC , klappen wir um die Achsen MB bzw. MC in die Zeichenebene.

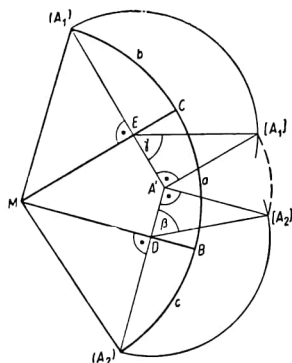


Abb. 16b

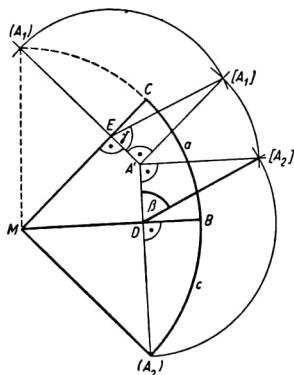


Abb. 16c

Die Stützdreiecke ADA' und AEA' werden um die Achsen DA' bzw. EA' umgelegt (Abb. 16b). Es ist $E(A_1) = E[A_1]$, $A'[A_1] = A'[A_2]$ und $D(A_2) = D[A_2]$.

Die Abbildung 16b zeigt, wie man konstruktiv die Größe der Winkel β und γ ermitteln kann, wenn die drei Seiten eines sphärischen Dreiecks gegeben sind. Durch Wahl einer anderen Seitenfläche als Grundriß lassen sich auch die Winkel α und γ bzw. α und β konstruieren, wobei durch die zweimal konstruierten Winkel eine gute Kontrolle über die Genauigkeit der Winkelwerte möglich ist.

Betrachtet man in der Abbildung 16a zwei Seiten und den eingeschlossenen Winkel (zum Beispiel α, β, c) als gegeben, so kann man entsprechend der Abbildung 16b durch Abänderung der Reihenfolge der Konstruktionen (vgl. Abb. 16c) die dritte fehlende Seite konstruieren. Man geht dabei von den Punkten $M, (A_2), B, C, D, A'$ und $[A_2]$ aus und verwendet die Beziehung $A'[A_2] = A'[A_1]$. Damit ist diese Aufgabe auf den Fall dreier gegebener Seiten zurückgeführt.

Alle durch Kanten begrenzten Körper bilden mehrere körperliche Ecken. Zum Beispiel hat der Würfel acht, das Tetraeder vier Ecken. Mit Hilfe der sphärischen Trigonometrie können die Neigungswinkel der Seitenflächen kantig begrenzter Körper berechnet werden, indem man die von den Kanten gebildeten Winkel als Seiten eines sphärischen n -Ecks auffaßt und die Neigungswinkel der Flächen als Winkel in den entsprechenden sphärischen n -Ecken berechnet.

So ist zum Beispiel eine Tetraederecke eine körperliche Ecke mit drei gleichen Seiten (Abb. 17). Ihre Seitenlänge ist 60° . Es muß jedoch darauf hingewiesen werden, daß nicht jede körperliche Ecke mit drei gleichen Seiten diese Seitenlänge

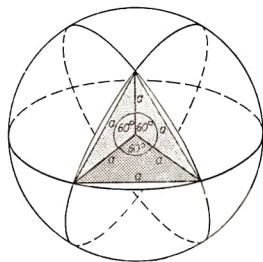


Abb. 17

hat. Durch einen Diagonalschnitt bei einer Oktaederecke erhält man eine körperliche Ecke mit zwei Seiten zu 60° und einer Seite zu 90° (Abb. 18a). Abbildung 18b zeigt die wahre Größe der Seitenflächen. Sie sind in die Zeichenebene umgeklappt.

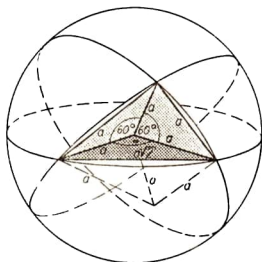


Abb. 18a

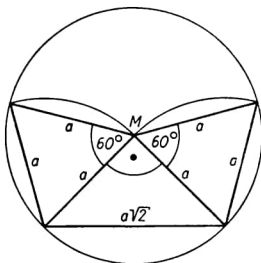


Abb. 18b

5. Die Polarecke und das Polardreieck, die Seitensumme

Es sei M der Scheitelpunkt einer körperlichen Ecke mit dem zugehörigen sphärischen Dreieck ABC . Im Scheitelpunkt M errichten wir auf den das Dreieck erzeugenden Ebenen MAB , MBC und MAC die Senkrechten. Diese bilden dort eine neue Ecke, die wir die **Polarecke** der ursprünglichen Ecke nennen. Die drei Senkrechten durchstoßen die Kugelfläche in den drei Punkten A' , B' und C' (Abb. 19a). Die Punkte A' , B' und C' bilden ein sphärisches Dreieck, das wir das **Polardreieck** $A'B'C'$ zu dem ursprünglichen Dreieck ABC nennen.

Seiten und Winkel des Polardreiecks werden genauso definiert wie Seiten und Winkel des sphärischen Dreiecks. Wir wollen nun an Hand der Abbildungen 19a und 19b die Zusammenhänge zwischen dem sphärischen Dreieck und seinem Polardreieck herleiten.

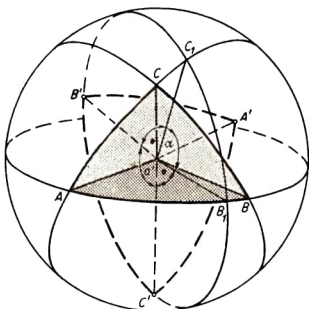


Abb. 19a

Die beiden Ebenen MAC und MAB haben den Neigungswinkel α , der gleich dem Winkel α des sphärischen Dreiecks bei A ist. Die Senkrechten im Punkt M auf AM in den Ebenen MAC und MAB haben bei M ebenfalls den Winkel α . Die beiden Schenkel dieses Winkels durchstoßen die Kugel in den Punkten C_1 und B_1 . Da die Kante MB' senkrecht auf der Ebene MAC steht, steht sie auch senkrecht auf MC_1 , und ebenso steht MC' senkrecht auf MB_1 . Die Punkte B' , C' , B_1 und C_1 liegen auf einem Großkreis, dessen wahre Größe in Abbildung 19b dargestellt wird. Die Seite $B'C'$

auf der Ebene MAC steht, steht sie auch senkrecht auf MC_1 , und ebenso steht MC' senkrecht auf MB_1 . Die Punkte B' , C' , B_1 und C_1 liegen auf einem Großkreis, dessen wahre Größe in Abbildung 19b dargestellt wird. Die Seite $B'C'$

des Polardreiecks nennen wir a' . Aus Abbildung 19b folgt die Beziehung

$$a + 90^\circ + a' + 90^\circ = 360^\circ$$

oder

$$a' = 180^\circ - a.$$

Die für die Seite a' durchgeführten Betrachtungen gelten entsprechend für die Seiten b' und c' . Somit ist

$$b' = 180^\circ - \beta \quad \text{und} \quad c' = 180^\circ - \gamma.$$

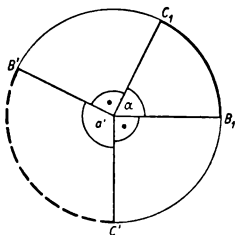


Abb. 19b

Wenn man das auf diese Weise gefundene Polardreieck $A'B'C'$ als Ausgangsdreieck wählt und dazu das Polardreieck konstruiert, erhält man wieder das Dreieck ABC . Demnach ist das ursprüngliche sphärische Dreieck das Polardreieck seines eigenen Polardreiecks. Da jedem Dreieck eine körperliche Ecke eindeutig zugeordnet ist, folgt daraus der Satz:

Jede Ecke ist die Polarecke ihrer Polarecke.

Aus diesem Satz ergeben sich drei weitere Formeln:

$$a = 180^\circ - a', \quad b = 180^\circ - \beta' \quad \text{und} \quad c = 180^\circ - \gamma'.$$

Diese sechs Formeln werden in einem Satz zusammengefaßt:

Die Seiten (Winkel) eines sphärischen Dreiecks ergänzen sich mit den entsprechenden Winkeln (Seiten) des Polardreiecks zu je 180° .

Aus diesem Satz folgt unter Verwendung der Sätze über die Winkelsumme im sphärischen Dreieck ein wichtiger Satz über die Summe der Seiten.

Es ist

$$180^\circ < a' + \beta' + \gamma'.$$

Unter Verwendung der obigen Formeln folgt dann

$$180^\circ < (180^\circ - a) + (180^\circ - b) + (180^\circ - c)$$

$$180^\circ < 540^\circ - (a + b + c)$$

$$a + b + c < 360^\circ.$$

Da a , b und c positiv vorausgesetzt sind, gilt also

$$0^\circ < a + b + c < 360^\circ.$$

Die Seitensumme im sphärischen Dreieck ist kleiner als 360° , das heißt kleiner als ein Großkreisumfang.

II. Berechnung des sphärischen Dreiecks

6. Das rechtwinklige sphärische Dreieck, die Nepersche Regel

Zur Ableitung der Formeln verwenden wir eine bei C rechtwinklige Ecke (Abb. 20a). Von A fallen wir die Lote AA' auf MC und AD auf MB und verbinden A' mit D . Die Ebene MCA steht senkrecht auf der Ebene MCB , also ist der Winkel $AA'D$ ein rechter Winkel. Der Winkel ADA' ist der Neigungswinkel der Ebene MBA gegen die Ebene MBC . Er ist der Winkel β , den die sphärischen Seiten $AB = c$ und $BC = a$ miteinander bilden. Außerdem ist $AC = b$.

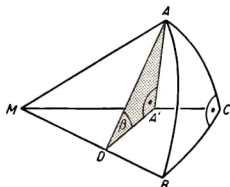


Abb. 20 a

Die Ebene MCB verwenden wir als Zeichenebene; wir legen die Ebenen MCA bzw. MBA um die Achsen MC bzw. MB um. Das Dreieck $AA'D$ legen wir um die Achse DA' um (Abb. 20b).

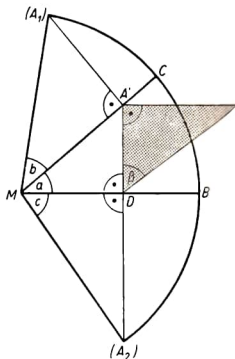


Abb. 20 b

Aus Abbildung 20b entnehmen wir die folgenden Beziehungen:

$$\begin{aligned} \sin a &= \frac{A'D}{MA'}, & \sin b &= \frac{A'(A_1)}{M(A_1)} = \frac{A'(A_1)}{r}, & \sin c &= \frac{D(A_2)}{M(A_2)} = \frac{D(A_2)}{r}, \\ \cos a &= \frac{MD}{MA'}, & \cos b &= \frac{MA'}{M(A_1)} = \frac{MA'}{r}, & \cos c &= \frac{MD}{M(A_2)} = \frac{MD}{r}, \\ \sin \beta &= \frac{A'[A]}{D[A]}, & \cos \beta &= \frac{A'D}{D[A]}. \end{aligned}$$

Wir bilden nun:

$$\cos a \cdot \cos b = \frac{MD}{MA'} \cdot \frac{MA'}{r} = \frac{MD}{r} = \cos c,$$

$$\cos a \cdot \cos b = \cos c; \quad (1)$$

$$\frac{\sin b}{\sin c} = \frac{A'(A_1)}{r} \cdot \frac{r}{D(A_2)} = \frac{A'(A_1)}{D(A_2)} = \frac{A'[A]}{D[A]} = \sin \beta,$$

$$\frac{\sin b}{\sin c} = \sin \beta. \quad (2a)$$

Durch Vertauschen der Katheten a und b sowie entsprechendes Vertauschen der Winkel a und β erhalten wir:

$$\frac{\sin a}{\sin c} = \sin \alpha. \quad (2b)$$

Weiter erhalten wir:

$$\cos b \cdot \sin \alpha = \frac{MA'}{r} \cdot \frac{\sin a}{\sin c} = \frac{MA'}{r} \cdot \frac{A'D}{MA'} \cdot \frac{r}{D(A_2)} = \frac{A'D}{D(A_2)} = \frac{A'D}{D[A]} = \cos \beta,$$

$$\cos b \cdot \sin \alpha = \cos \beta \quad (3a)$$

und entsprechend durch Vertauschen der Katheten und Winkel

$$\cos a \cdot \sin \beta = \cos \alpha. \quad (3b)$$

Daraus folgt:

$$\cos a = \frac{\cos \alpha}{\sin \beta} \quad \text{und} \quad \cos b = \frac{\cos \beta}{\sin \alpha}.$$

Also wird aus (1)

$$\cos c = \cos a \cdot \cos b = \frac{\cos \alpha}{\sin \beta} \cdot \frac{\cos \beta}{\sin \alpha} = \operatorname{ctg} \alpha \cdot \operatorname{ctg} \beta,$$

$$\cos c = \operatorname{ctg} \alpha \cdot \operatorname{ctg} \beta. \quad (4)$$

Aus (3a) folgt:

$$\sin \alpha = \frac{\cos \beta}{\cos b},$$

aus (2a) folgt:

$$\sin c = \frac{\sin b}{\sin \beta},$$

also wird aus (2b):

$$\sin a = \sin \alpha \cdot \sin c = \frac{\cos \beta}{\cos b} \cdot \frac{\sin b}{\sin \beta} = \operatorname{ctg} \beta \cdot \operatorname{tg} b,$$

$$\sin a = \operatorname{ctg} \beta \cdot \operatorname{tg} b \quad (5a)$$

und entsprechend durch Vertauschen

$$\sin b = \operatorname{tg} a \cdot \operatorname{ctg} \alpha. \quad (5b)$$

Aus (1) folgt:

$$\cos a = \frac{\cos c}{\cos b},$$

nach (2a) ist

$$\sin \beta = \frac{\sin b}{\sin c},$$

also wird aus (3b)

$$\cos \alpha = \cos a \cdot \sin \beta = \frac{\cos c}{\cos b} \cdot \frac{\sin b}{\sin c} = \operatorname{tg} b \cdot \operatorname{ctg} c,$$

$$\cos \alpha = \operatorname{tg} b \cdot \operatorname{ctg} c \quad (6a)$$

und entsprechend durch Vertauschen

$$\cos \beta = \operatorname{tg} a \cdot \operatorname{ctg} c. \quad (6b)$$

Damit haben wir zehn Formeln zur Berechnung des rechtwinkligen sphärischen Dreiecks erhalten. Sie lassen sich in einer Regel zusammenfassen.

Ordnet man die Bestimmungsstücke eines rechtwinkligen sphärischen Dreiecks wie in Abbildung 21, wobei man den rechten Winkel ausläßt und die Katheten a und b durch ihre Komplemente $(90^\circ - a)$ und $(90^\circ - b)$ ersetzt, so ist der Kosinus eines jeden Stückes gleich

- dem Produkt der Kotangenten der anliegenden Stücke,
- dem Produkt der Sinus der gegenüberliegenden Stücke.

Diese Regel heißt **Nepersche Regel**¹⁾.

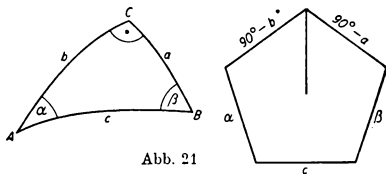


Abb. 21

7. Das schiefwinklige Dreieck

Nunmehr wollen wir Formeln zur Berechnung eines schiefwinkligen sphärischen Dreiecks herleiten.

a) Der Sinussatz der sphärischen Trigonometrie

Wir gehen wie in der ebenen Trigonometrie vor und teilen durch eine Höhe das gegebene Dreieck in zwei rechtwinklige Teildreiecke, auf die wir die Nepersche Regel anwenden können.

Die Abbildung 22 ist die Skizze eines sphärischen Dreiecks mit den zugehörigen Neperschen Fünfecken für die beiden rechtwinkligen Teildreiecke. Es ist in ADC

$$\sin h_c = \sin \alpha \cdot \sin b$$

und in BDC

$$\sin h_c = \sin \beta \cdot \sin a.$$

Daraus folgt

$$\sin \alpha \cdot \sin b = \sin \beta \cdot \sin a$$

oder

$$\frac{\sin \alpha}{\sin b} = \frac{\sin \alpha}{\sin \beta}.$$

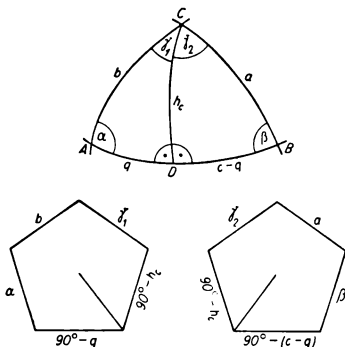


Abb. 22

¹⁾ Diese Regel wird nach dem englischen Mathematiker John Napier (1550 bis 1617; latinisierte Form des Namens; Neper) benannt, der 1614 eine ähnliche Regel gleichzeitig mit seiner Logarithmentafel veröffentlichte.

Durch zyklische Vertauschung oder durch Verwendung der anderen beiden Höhen h_a und h_b lassen sich dazu noch die beiden Formeln

$$\frac{\sin b}{\sin c} = \frac{\sin \beta}{\sin \gamma} \quad \text{und} \quad \frac{\sin a}{\sin c} = \frac{\sin \alpha}{\sin \gamma}$$

ableiten.

Diese drei Formeln, die eine gewisse Ähnlichkeit mit denen der ebenen Trigonometrie haben, lassen sich auch kurz als laufende Proportion

$$\sin \alpha : \sin \beta : \sin \gamma = \sin a : \sin b : \sin c,$$

oder in der Form

$$\frac{\sin a}{\sin \alpha} = \frac{\sin b}{\sin \beta} = \frac{\sin c}{\sin \gamma}$$

schreiben.

Diese Gleichungen werden als der **Sinussatz der sphärischen Trigonometrie** bezeichnet.

b) Die Kosinussätze der sphärischen Trigonometrie

1. Der Seitenkosinussatz

Wir wollen nun einen dem Kosinussatz der ebenen Trigonometrie entsprechenden Satz für die sphärische Trigonometrie herleiten. Die dazu erforderlichen Beziehungen werden unter Verwendung der Neperschen Regel aus den beiden Teildreiecken gewonnen.

Es gilt im Dreieck BDC (Abb. 22)

$$\cos a = \cos h_c \cdot \cos (c - q), \quad (\text{I})$$

im Dreieck ADC

$$\cos b = \cos h_c \cdot \cos q, \quad (\text{II})$$

$$\sin q = \operatorname{ctg} \alpha \cdot \operatorname{tg} h_c = \frac{\cos \alpha}{\sin \alpha} \cdot \frac{\sin h_c}{\cos h_c} \quad (\text{III})$$

und

$$\sin h_c = \sin a \cdot \sin b. \quad (\text{IV})$$

Unter Verwendung eines Additionstheorems geht die Formel (I) über in

$$\cos a = \cos h_c \cdot \cos c \cdot \cos q + \cos h_c \cdot \sin c \cdot \sin q. \quad (\text{I}')$$

Aus Gleichung (III) folgt

$$\cos h_c \cdot \sin q = \frac{\cos \alpha \cdot \sin h_c}{\sin \alpha}$$

und durch Einsetzen von (IV)

$$\cos h_c \cdot \sin q = \frac{\cos \alpha \cdot \sin a \cdot \sin b}{\sin \alpha} = \sin b \cdot \cos \alpha. \quad (\text{III}')$$

Durch Einsetzen von (II) und (III') in (I') folgt schließlich

$$\cos a = \cos b \cdot \cos c + \sin b \cdot \sin c \cdot \cos \alpha.$$

Durch zyklische Vertauschung oder durch Verwendung der beiden anderen Höhen folgen die Formeln

$$\cos b = \cos a \cdot \cos c + \sin a \cdot \sin c \cdot \cos \beta$$

und

$$\cos c = \cos a \cdot \cos b + \sin a \cdot \sin b \cdot \cos \gamma.$$

Diese drei Formeln, die neben den drei Seiten jeweils nur einen Winkel enthalten, bilden den **Seitenkosinussatz der sphärischen Trigonometrie**.

2. Der Winkelkosinussatz

Die Beziehungen zwischen den Seiten und Winkeln eines sphärischen Dreiecks und seines Polardreiecks ermöglichen es, aus dem Seitenkosinussatz weitere Formeln abzuleiten. In

$$\cos a' = \cos b' \cdot \cos c' + \sin b' \cdot \sin c' \cdot \cos a',$$

dem auf das Polardreieck $A'B'C'$ angewandten Seitenkosinussatz, setzt man:

$$a' = 180^\circ - \alpha, \quad b' = 180^\circ - \beta, \quad c' = 180^\circ - \gamma \quad \text{und} \quad a' = 180^\circ - a.$$

Man erhält zunächst

$$\begin{aligned} \cos (180^\circ - \alpha) &= \cos (180^\circ - \beta) \cdot \cos (180^\circ - \gamma) \\ &\quad + \sin (180^\circ - \beta) \cdot \sin (180^\circ - \gamma) \cdot \cos (180^\circ - \alpha). \end{aligned}$$

Daraus folgt

$$-\cos \alpha = \cos \beta \cdot \cos \gamma - \sin \beta \cdot \sin \gamma \cdot \cos a$$

und nach Multiplikation mit (-1)

$$\cos \alpha = -\cos \beta \cdot \cos \gamma + \sin \beta \cdot \sin \gamma \cdot \cos a.$$

Durch zyklische Vertauschung folgen die weiteren Formeln

$$\cos \beta = -\cos \alpha \cdot \cos \gamma + \sin \alpha \cdot \sin \gamma \cdot \cos b$$

und

$$\cos \gamma = -\cos \alpha \cdot \cos \beta + \sin \alpha \cdot \sin \beta \cdot \cos c.$$

Diese drei Formeln, die neben den drei Winkeln jeweils nur eine Seite enthalten, bilden den **Winkelkosinussatz der sphärischen Trigonometrie**.

Bei der Ableitung dieses Satzes haben wir die Beziehungen zwischen dem sphärischen Dreieck und seinem Polardreieck verwendet. Wir sagen deshalb, der Seitenkosinussatz und der Winkelkosinussatz sind einander polar.

Man kann ähnliche Überlegungen auch für den Sinussatz durchführen, erhält jedoch keine neuen Formeln. Wir sagen deshalb, der Sinussatz ist zu sich selbst polar.

8. Die sechs Fälle der Dreiecksberechnung

Ein ebenes Dreieck ist durch die drei Winkel nicht eindeutig bestimmt, da durch die Winkelsummenbeziehung der dritte Winkel aus zwei gegebenen Winkeln berechnet werden kann. Das gilt für das sphärische Dreieck nicht mehr. Durch die drei Winkel eines sphärischen Dreiecks sind die drei Seiten des zugehörigen Polardreiecks eindeutig bestimmt. Daraus folgt, daß das Polardreieck und damit das ursprüngliche sphärische Dreieck eindeutig bestimmt sind. Es sind also die folgenden sechs Fälle der Dreiecksberechnung oder Dreieckskonstruktion möglich (zur Abkürzung bezeichnen wir eine Seite mit S und einen Winkel mit W):

$S S S$ (3 Seiten),

$S W S$ (zwei Seiten und der eingeschlossene Winkel),

$S S W$ (zwei Seiten und ein Gegenwinkel)

und $W W W$ (drei Winkel),

$W S W$ (zwei Winkel und die Zwischenseite),

$W W S$ (zwei Winkel und eine Gegenseite).

Die drei Fälle der zweiten Gruppe lassen sich entsprechend den Ausführungen des ersten Absatzes auf die drei Fälle der ersten Gruppe zurückführen.

Die Konstruktion der Fälle SSS und SWS ist im Abschnitt 4 in den Abbildungen 16b und 16c dargestellt. Die Konstruktion der Fälle WWW und WSW wird mit Hilfe des Polardreiecks auf die Fälle SSS und SWS zurückgeführt. Die Berechnung dieser vier Fälle bietet keine Schwierigkeiten. Der Fall SSS kann durch dreimalige Verwendung des Seitenkosinussatzes und der Fall WWW entsprechend durch dreimalige Verwendung des Winkelkosinussatzes gelöst werden. Da für Eulersche Dreiecke nur Seiten bzw. Winkel unter 180° in Frage kommen, ist durch Verwendung der Kosinussätze die Lösung eindeutig bestimmt, denn die Kosinusfunktion hat im ersten und zweiten Quadranten verschiedene Vorzeichen.

In den Fällen SSS und WWW genügt es auch, wenn man nur einen Winkel bzw. eine Seite mit einem Kosinussatz berechnet, um dann die fehlenden beiden Stücke mit Hilfe des Sinussatzes zu ermitteln. Da aber die Sinusfunktion im ersten und zweiten Quadranten gleiche Vorzeichen hat, ist diese Lösung nicht eindeutig. Diesen Nachteil muß man jedesmal bedenken, wenn man eine Berechnung mit Hilfe des Sinussatzes vornimmt. Es gibt jedoch wie bei den ebenen Dreiecken Beziehungen zwischen den Winkeln und den Seiten eines sphärischen Dreiecks, die es häufig ermöglichen, von den beiden möglichen Lösungen die brauchbare zu bestimmen. Die wichtigsten Sätze seien hier ohne Beweis erwähnt.

1. Im sphärischen Dreieck liegen der kleineren Seite der kleinere Winkel, der mittleren Seite der mittlere Winkel, der größeren Seite der größere Winkel und gleichen Seiten gleiche Winkel gegenüber.

2. Die Summe zweier Seiten ist stets größer als die dritte Seite.
3. Die Summe zweier Seiten und die Summe der beiden Gegenwinkel sind stets zugleich kleiner, gleich oder größer als 180° .
4. Die Summe zweier Winkel ist stets kleiner als der um 180° vermehrte dritte Winkel.

In den Fällen *SWS* und *WSW* erhält man bei der Anwendung des Kosinussatzes die fehlende dritte Seite bzw. den fehlenden dritten Winkel. Sie sind damit auf die Fälle *SSS* und *WWW* zurückgeführt. In den Fällen *SSW* und *WWS* findet man zunächst unter Verwendung des Sinussatzes den fehlenden gegenüberliegenden Winkel bzw. die fehlende gegenüberliegende Seite. Dann bereitet aber die weitere Lösung Schwierigkeiten, weil die beiden noch fehlenden Stücke selbst einander gegenüberliegen, so daß weder der Sinussatz noch einer der Kosinussätze unmittelbar zum Ziel führen.

Wir wollen diese Fälle an einem konkreten Beispiel behandeln. Für den Fall *SSW* seien a , b und α gegeben. Der Winkel β wird sofort mit Hilfe des Sinussatzes gefunden:

$$\sin \beta = \frac{\sin b \cdot \sin \alpha}{\sin a}.$$

Die unbekannte Seite c und der unbekannte Winkel γ werden mit Hilfe der Höhe h_c zerlegt. Es entstehen zwei rechtwinklige Teildreiecke, aus denen mit Hilfe der Neperschen Formeln γ_1 , γ_2 , c_1 und c_2 berechnet werden können. Je nachdem nun α ein spitzer oder ein stumpfer Winkel ist, wird $\gamma = \gamma_2 \pm \gamma_1$ und $c = c_2 \pm c_1$ (vgl. Abb. 23).

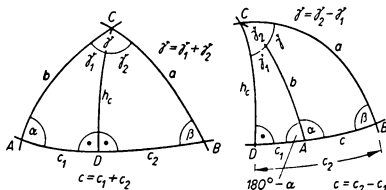


Abb. 23

Einen zweiten, aber rechnerisch umständlichen Lösungsweg erhält man, indem man auf die noch unbekannte Seite und den noch unbekannten Winkel den Seiten- bzw. Winkelkosinussatz anwendet. Es ist also

$$\cos c = \cos a \cdot \cos b + \sin a \cdot \sin b \cdot \cos \gamma$$

und

$$\cos \gamma = -\cos a \cdot \cos \beta + \sin a \cdot \sin \beta \cdot \cos c.$$

Es ergibt sich also ein Gleichungssystem mit den beiden Unbekannten $\cos c$ und $\cos \gamma$.

Bei einem dritten, häufig begangenen Lösungsweg verwendet man einen Kosinussatz, der auf eine schon bekannte Seite führt. Die unbekannte Größe tritt dann gleichzeitig in der Kosinusfunktion und in der Sinusfunktion auf. Es ist zum Beispiel

$$\cos a = \cos b \cdot \cos c + \sin b \cdot \sin c \cdot \cos \alpha. \quad (\text{I})$$

Wir verwenden nun als Hilfsgröße den Bogen AD und bezeichnen ihn mit ψ (vgl. Abb. 23). Dann ist nach der Neperschen Regel

$$\cos a = \operatorname{ctg} (90^\circ - \psi) \cdot \operatorname{ctg} b,$$

$$\cos a = \operatorname{tg} \psi \cdot \operatorname{ctg} b,$$

mithin

$$\frac{\cos a}{\operatorname{ctg} b} = \operatorname{tg} \psi$$

oder

$$\cos a \cdot \operatorname{tg} b = \operatorname{tg} \psi; \quad (\text{II})$$

daraus erhalten wir

$$\cos a \cdot \sin b = \cos b \cdot \operatorname{tg} \psi. \quad (\text{III})$$

Diesen Ausdruck setzen wir in die Gleichung (I) ein. Dann erhalten wir

$$\cos a = \cos b \cdot \cos c + \sin c \cdot \cos b \cdot \operatorname{tg} \psi.$$

Durch weitere Umformung folgt dann

$$\cos \psi \cdot \cos a = \cos b \cdot \cos c \cdot \cos \psi + \sin c \cdot \cos b \cdot \sin \psi,$$

$$\cos \psi \cdot \cos a = \cos b \cdot (\cos c \cdot \cos \psi + \sin c \cdot \sin \psi).$$

Wir wenden auf diese Formel ein Additionstheorem an und erhalten

$$\cos \psi \cdot \cos a = \cos b \cdot \cos (c - \psi).$$

Daraus folgt

$$\cos (c - \psi) = \frac{\cos \psi \cdot \cos a}{\cos b},$$

und da ψ aus der Gleichung (II) bestimmt ist, ist auch c bestimmbar.

Aufgaben

- Wie lauten die 10 Formeln für das rechtwinklige sphärische Dreieck, wenn das Dreieck a) zwei rechte Winkel, b) drei rechte Winkel hat?
- Berechnen Sie die fehlenden Stücke des rechtwinkligen sphärischen Dreiecks (Hypotenuse c)!

a) $a = 55,2^\circ$, $b = 47,5^\circ$	b) $a = 91,9^\circ$, $\beta = 42,3^\circ$
c) $a = 48,5^\circ$, $c = 97,2^\circ$	d) $a = 83,6^\circ$, $a = 71,1^\circ$
e) $b = 135,6^\circ$, $\beta = 99,0^\circ$	i) $a = 99,0^\circ$, $\alpha = 135,6^\circ$
g) $\alpha = 132,2^\circ$, $\beta = 104,6^\circ$	h) $b = 167,4^\circ$, $\alpha = 73,9^\circ$
j) $b = 103,1^\circ$, $c = 92,6^\circ$	k) $\alpha = 117,3^\circ$, $\beta = 32,5^\circ$

Wieviel Lösungen gibt es jeweils? Beachten Sie insbesondere e) und f)!
- Berechnen Sie die fehlenden Stücke des gleichschenkligen sphärischen Dreiecks durch Zurückführung auf zwei rechtwinklige Dreiecke (Basis b)!

a) $a = 117,9^\circ$, $b = 35,6^\circ$	b) $b = 97,2^\circ$, $\alpha = 84,4^\circ$
c) $a = 87,2^\circ$, $\beta = 119,3^\circ$	d) $a = 78,4^\circ$, $\alpha = 57,7^\circ$

4. Unter welchem Winkel sind zwei benachbarte Seitenflächen bei einem
 a) Tetraeder, b) Oktaeder, c) Dodekaeder und d) Ikosaeder
 gegeneinander geneigt (vgl. Abb. 17 und 18)?
5. Unter welchem Winkel sind die Seitenflächen einer geraden, regelmäßigen, vierseitigen Pyramide gegeneinander geneigt, wenn je zwei benachbarte Kanten an der Spitze einen Winkel von a) 20° , b) 40° , c) 60° , d) 80° bilden?
6. In den folgenden Aufgaben sind die fehlenden Seiten und Winkel der sphärischen Dreiecke zu berechnen.
- | | |
|--|---|
| a) $a = 21,3^\circ$, $b = 33,6^\circ$, $c = 39,5^\circ$ | b) $a = 114,3^\circ$, $b = 89,9^\circ$, $c = 121,5^\circ$ |
| c) $a = 93,4^\circ$, $b = 107,3^\circ$, $c = 59,5^\circ$ | d) $a = 74,7^\circ$, $b = 98,6^\circ$, $c = 65,6^\circ$ |
| e) $\alpha = 83,4^\circ$, $\beta = 63,3^\circ$, $\gamma = 79,9^\circ$ | f) $\alpha = 65,6^\circ$, $\beta = 112,6^\circ$, $\gamma = 118,7^\circ$ |
| g) $\alpha = 143,0^\circ$, $\beta = 167,8^\circ$, $\gamma = 152,3^\circ$ | h) $\alpha = 84,5^\circ$, $\beta = 100,9^\circ$, $\gamma = 70,6^\circ$ |
| i) $c = 51,5^\circ$, $b = 42,9^\circ$, $\gamma = 48,0^\circ$ | k) $c = 127,3^\circ$, $a = 42,5^\circ$, $\gamma = 45,7^\circ$ |
| l) $a = 128,5^\circ$, $b = 89,2^\circ$, $\alpha = 135,7^\circ$ | m) $a = 131,3^\circ$, $c = 119,9^\circ$, $\alpha = 72,2^\circ$ |
| n) $a = 46,5^\circ$, $\beta = 73,9^\circ$, $\alpha = 84,4^\circ$ | o) $\alpha = 114,4^\circ$, $\gamma = 63,6^\circ$, $a = 55,2^\circ$ |
| p) $\beta = 61,1^\circ$, $\gamma = 119,8^\circ$, $b = 126,5^\circ$ | q) $\beta = 79,4^\circ$, $\alpha = 59,6^\circ$, $c = 103,3^\circ$ |
| r) $a = 97,7^\circ$, $c = 119,9^\circ$, $\beta = 140,5^\circ$ | s) $a = 89,4^\circ$, $b = 33,8^\circ$, $\gamma = 178,2^\circ$ |
| t) $b = 108,6^\circ$, $c = 107,7^\circ$, $\alpha = 108,6^\circ$ | u) $a = 104,5^\circ$, $b = 42,3^\circ$, $\gamma = 79,6^\circ$ |
| v) $\alpha = 107,2^\circ$, $\beta = 33,3^\circ$, $c = 87,9^\circ$ | w) $\beta = 129,2^\circ$, $\gamma = 109,1^\circ$, $a = 85,3^\circ$ |
| x) $\alpha = 83,2^\circ$, $\gamma = 24,8^\circ$, $b = 38,5^\circ$ | y) $\alpha = 49,3^\circ$, $\beta = 138,7^\circ$, $c = 12,8^\circ$ |

Bemerkung: Prüfen Sie nach Lösung der Aufgaben, ob in jedem Falle der Satz 1 (S. 91) erfüllt ist! Aufgabe s sorgfältig interpolieren!

7. Berechnen Sie den Winkel in einem gleichseitigen sphärischen Dreieck, wenn die Seite
 a) $\alpha = 10^\circ$, b) $\alpha = 30^\circ$, c) $\alpha = 50^\circ$, d) $\alpha = 70^\circ$, e) $\alpha = 90^\circ$, f) $\alpha = 110^\circ$ gegeben ist!
8. Es ist die Seite in einem gleichseitigen sphärischen Dreieck zu errechnen, wenn der Winkel
 a) $\alpha = 70^\circ$, b) $\alpha = 90^\circ$, c) $\alpha = 110^\circ$, d) $\alpha = 130^\circ$, e) $\alpha = 150^\circ$, f) $\alpha = 170^\circ$ gegeben ist!
9. Zwischen welchen Grenzen liegt der Winkel in einem gleichseitigen sphärischen Dreieck? Berechnen Sie die zu diesen Grenzen gehörigen Seiten!
10. Die Kanten eines Dreibeins betragen
 $r = 49$ cm, $s = 37$ cm, $t = 42$ cm;
 sie bilden die Winkel
 $(r, s) = \alpha = 80^\circ$, $(s, t) = \beta = 45^\circ$,
 $(t, r) = \gamma = 60^\circ$.
 Unter welchem Winkel sind die
 durch die Kanten bestimmten Ebenen
 gegen die waagerechte Ebene
 geneigt, auf die das Dreibein ge-
 stellt wird (vgl. Abb. 24)?

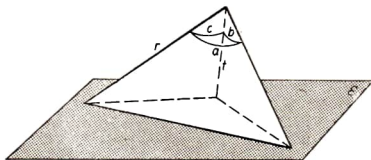


Abb. 24

III. Anwendung auf die Erdkugel

9. Das Gradnetz und die Größe der Erde

Durch neuere Messungen sowjetischer Geodäten wurde bestätigt, daß die Erde ein Ellipsoid mit drei ungleichen Achsen ist. Die lineare Exzentrizität des Äquators ist jedoch so klein, daß sie bei praktischen Berechnungen keine Rolle spielt. Man kann also den Äquator als Kreis und mithin die Erde als ein Ellipsoid mit zwei ungleichen Achsen annehmen. Die große Halbachse des Ellipsoides (Äquatorradius) hat die Länge von rund 6378 km, die kleine Halbachse (Erdmittelpunkt . . . Pol) die Länge von rund 6357 km. In vielen Fällen kann man sogar die Abweichung des Rotationsellipsoides von der Kugel vernachlässigen. Es ist dann üblich, mit dem mittleren Erdradius von $r = 6370$ km zu rechnen. Die dadurch auftretenden Fehler sind im allgemeinen unwesentlich. Wir wollen im folgenden für alle Berechnungen diesen Wert zugrunde legen, soweit nichts anderes angegeben ist. Durch diese Vereinfachung wird es uns möglich, die Ergebnisse der sphärischen Trigonometrie auf die Erde anzuwenden.

Aus dem Erdkundeunterricht ist uns das Koordinatensystem auf der Erde bekannt. Durch dieses Koordinatensystem, das Gradnetz der Erde, ist die Lage jedes Ortes P auf der Erde eindeutig angebar. Wir wollen im folgenden eine kurze Zusammenfassung der wichtigsten Tatsachen geben.

Durch jeden Ort P und die beiden Pole N und S (geographischer Nordpol und geographischer Südpol) ist ein Großkreis bestimmt. Der Halbkreis von einem Pol über den Ort P zum anderen Pol wird als Meridian dieses Ortes bezeichnet. Der Meridian von Greenwich ist der Nullmeridian. Der sphärische Winkel des Kugelzweiecks, das durch den Nullmeridian und den Meridian des Ortes P gebildet wird, ist die **geographische Länge** des Ortes (Abb. 25). Dabei wird der Winkel gemessen, der kleiner als 180° ist; die

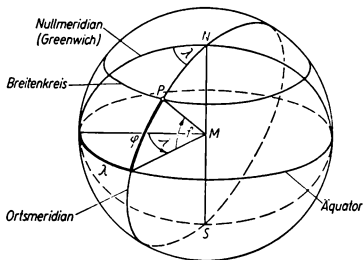


Abb. 25

Länge wird als östliche oder westliche Länge angegeben, je nachdem, ob die Messung vom Nullmeridian aus in östlicher oder in westlicher Richtung erfolgt. Die östliche ist dabei die positive und die westliche die negative Richtung.

Der Durchmesser durch die beiden Pole ist die Erdachse. Der Großkreis, auf dessen erzeugender Ebene die Erdachse senkrecht steht, heißt Äquator. Durch jeden Ort P , der nicht auf dem Äquator liegt, wird ein Kleinkreis bestimmt, der dem Äquator parallel ist. Dieser Kleinkreis heißt **Breitenkreis** (oder — in der Nautik — Breitenparallel). Durch den Äquator und den Breitenkreis wird auf dem Meridian von P ein Bogen ausgeschnitten. Der Winkel, der zu diesem Bogen gehört, ist die **geographische Breite** des Ortes (Abb. 25). Die Breite des Ortes P heißt nördliche Breite, wenn vom Äquator aus nach Norden gemessen wird, und südliche Breite,

wenn vom Äquator aus nach Süden gemessen wird. Der nördlichen Breite gibt man auch das positive Vorzeichen, der südlichen entsprechend das negative.

Bei Aufgaben, in denen die Entfernung zweier Orte auf der Erdoberfläche bestimmt werden soll, ist es notwendig, das Ergebnis in einem Längenmaß anzugeben. Die Berechnung der Entfernung zweier Orte ergibt zunächst nur die Länge des Bogens im Gradmaß. Für die Umrechnung in Längeneinheiten wurden Mittelwerte für einen Meridiangrad bzw. einen Äquatorgrad bestimmt. In der Nautik wurde darüber hinaus eine besondere Längeneinheit, die Seemeile, eingeführt.

Eine Seemeile ist die Länge einer mittleren Meridianminute: $1' \triangleq 1852 \text{ m}$. Diese Beziehung beruht auf den Messungen, die zur Festlegung der Längeneinheit Meter durchgeführt wurden. Danach wurde das Meter als der zehnmillionste Teil des Erdquadranten definiert. Für die Länge eines Meridiangrades ergibt sich demnach $111,1 \text{ km}$. Diese Werte für die Längeneinheiten weichen nur sehr wenig von den Werten ab, die sich ergeben, wenn für die Berechnung der Längeneinheiten der Erdradius mit $r = 6370 \text{ km}$ angenommen wird. Eine weitere Längeneinheit, die durch die Länge eines Bogens definiert wurde, ist die geographische Meile. Eine geographische Meile ist die Länge von vier Äquatorminuten: $1 \text{ geographische Meile} = 7420 \text{ m}$. Die Länge eines Äquatorgrades beträgt $111,3 \text{ km}$. Als Maßeinheiten werden also die folgenden Werte verwendet:

$$1 \text{ mittlerer Meridiangrad} \triangleq 111,1 \text{ km} = 60 \text{ sm}$$

$$1 \text{ sm} = 1852 \text{ m},$$

$$1 \text{ Äquatorgrad} \triangleq 111,3 \text{ km} = 15 \text{ geographische Meilen}$$

$$1 \text{ geographische Meile} = 7420 \text{ m}.$$

Da die Ungenauigkeit für einen Meridiangrad höchstens $\pm 0,18\%$ und für einen Äquatorgrad höchstens $-0,36\%$ beträgt, liefern beide Werte bei Entfernungsberechnungen hinreichend gute Näherungen. In der Nautik ist es üblich, die Seemeile zu verwenden.

10. Das Poldreieck, die Orthodrome und die Kurswinkel

Zwei Orte A und B auf der Erde (Abb. 26a) sind durch ihre Koordinaten $A(\varphi_1; \lambda_1)$ und $B(\varphi_2; \lambda_2)$ gegeben. Der kleinere Bogen des Großkreises durch A und B ist die kürzeste Entfernung l zwischen den beiden Orten. Eine in der Nautik und der mathematischen Geographie häufig auftretende Aufgabe ist es, die kürzeste Entfernung zweier Orte auf der Erde zu berechnen. Zur Berechnung dieser Entfernung verwenden wir das sphärische Dreieck ANB . Dieses Dreieck heißt das **Poldreieck**, weil der Nordpol der dritte Eckpunkt N ist. Für zwei Orte auf der Südhalbkugel verwendet man zweckmäßig das Poldreieck, das den Südpol als dritten Eckpunkt hat. Liegt ein Ort auf der Nord- und der andere auf der Südhalbkugel, so ist es gleich, welchen der Pole man verwendet.

Wir zeichnen als Analysisfigur für das Poldreieck ANB ein durch drei Bogen begrenztes Dreieck (Abb. 26b). Dabei ist $AB = l$ die gesuchte kürzeste Entfernung, $AN = 90^\circ - \varphi_1$ und $BN = 90^\circ - \varphi_2$ jeweils das Komplement der geographischen Breite und der Winkel $ANB = \lambda_2 - \lambda_1 = \Delta\lambda$ der Längenunterschied der beiden Orte. Die kürzeste Entfernung l ist nach dem Seitenkosinussatz

$$\cos l = \cos (90^\circ - \varphi_1) \cdot \cos (90^\circ - \varphi_2) + \sin (90^\circ - \varphi_1) \cdot \sin (90^\circ - \varphi_2) \cdot \cos \Delta \lambda$$

oder nach Umformung

$$\cos l = \sin \varphi_1 \cdot \sin \varphi_2 + \cos \varphi_1 \cdot \cos \varphi_2 \cdot \cos \Delta \lambda.$$

Das Ergebnis ist eindeutig, da die Kosinusfunktion im ersten und zweiten Quadranten verschiedene Vorzeichen hat.

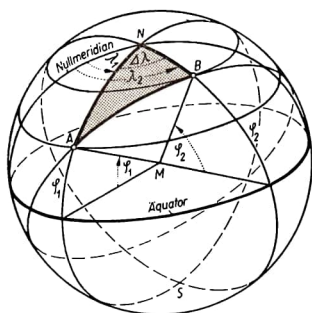


Abb. 26 a

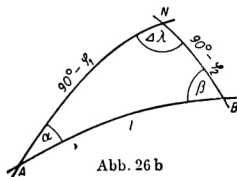


Abb. 26 b

Der Großkreisbogen zwischen *A* und *B* heißt auch **Orthodrome**, dementsprechend wird die kürzeste Entfernung auch **orthodrome Entfernung** genannt.

Von den beiden Tangenten im Punkt *A* an die Großkreise des Poldreiecks zeigt die eine Tangente nach Norden und die andere in die Himmelsrichtung, nach der

man sich wenden muß, wenn man auf der kürzesten Verbindung nach *B* gelangen will. Der Winkel α wird deshalb der **Kurswinkel des Anfangskurses** genannt. Eine ähnliche Überlegung gilt im Punkte *B*, nur daß hier die zweite Tangente in die Richtung weist, aus der man in *B* ankommt. Der Kurswinkel β heißt deshalb auch der **Kurswinkel des Endkurses**. Wollte man in *B* auf demselben Großkreis weiterfahren, so wäre der zugehörige Anfangskurswinkel $180^\circ - \beta$. Im allgemeinen, das heißt, wenn $\varphi_1 \neq \varphi_2$ und $\lambda_1 \neq \lambda_2$ ist, ist $\alpha \neq 180^\circ - \beta$. Daraus können wir schließen, daß der Kurswinkel auf der Orthodromen seinen Wert ändert. Diese Änderung erfolgt stetig. Aus diesem Grunde wird im allgemeinen nicht der orthodrome Kurs gesteuert.

Die Berechnung der beiden Kurswinkel erfolgt unter Anwendung des Sinussatzes. Die Werte sind nicht eindeutig und erfordern die Heranziehung der auf den Seiten 91 bis 92 erwähnten Sätze. Bei Verwendung des Seitenkosinussatzes lassen sich auch die beiden Kurswinkel in jedem Fall eindeutig bestimmen. Wegen der einfacheren Rechnung ist jedoch im allgemeinen die Anwendung des Sinussatzes vorzuziehen.

Die Bezeichnung der Kurswinkel ist nicht einheitlich. Im allgemeinen gibt man die Winkel nach der folgenden Methode an: Man bezeichnet die Nordrichtung mit 0° und weiter die Ostrichtung mit 90° , die Südrichtung mit 180° und die Westrichtung mit 270° (vgl. die Einteilung einer Windrose). Die Südostrichtung bildet also mit der Nordrichtung einen Winkel von 135° . Man schreibt sie: N 135° O (lies: Nord, 135° über Ost). In einzelnen Fällen wird noch die folgende Schreibweise durchgeführt: Man gibt der Nordrichtung und der Südrichtung die Bezeichnung 0°

und mißt nun den Winkel von 0° bis 90° in östlicher oder westlicher Richtung. Für die Südostrichtung ergibt sich danach die folgende Schreibweise: S 45° O (Süd, 45° über Ost).

11. Loxodrome, Scheitelpunkt, Schnitt mit Meridian- und Parallelkreisen

Da es im allgemeinen praktisch nicht möglich ist, auf der Orthodromen zu fahren, ist es notwendig, eine Linie konstanten Kurses festzulegen. Die Linie konstanten Kurses hat die Eigenschaft, daß sie sämtliche Meridiane unter dem gleichen Winkel schneidet und heißt **Loxodrome**.

Verlauf und Länge der Loxodrome lassen sich im allgemeinen nicht mit den Mitteln der sphärischen Trigonometrie berechnen. Die Bestimmung der Loxodrome ist ein Problem der Differentialgeometrie, also des Zweiges der Mathematik, der sich mit der Anwendung der Infinitesimalrechnung auf die Geometrie beschäftigt. Die Loxodrome verläuft im allgemeinen wellenförmig auf der Erdkugel und nähert sich asymptotisch, je nach der Fahrtrichtung, zwei gegenüberliegenden Polen. Von Gerhard Kremer (genannt Mercator, 1512 bis 1594) ist eine Karte entwickelt worden, auf der sämtliche Loxodromen gerade Linien sind. Diese Karte wird Mercatorkarte genannt (vgl. Atlas zur Erd- und Länderkunde, Große Ausgabe, Karten 92/94 c bis e). Damit ist es möglich, durch einfaches Anlegen eines Lineals an diese Karte den loxodromen Kurs abzulesen. Die loxodrome Entfernung zweier Orte ist nie kleiner als die orthodrome Entfernung. Liegen beide Orte auf dem gleichen Meridian oder auf dem Äquator, so sind beide Entfernungen gleich. In allen anderen Fällen ist die loxodrome Entfernung größer. Der Unterschied kann beträchtliche Ausmaße annehmen. So beträgt zum Beispiel der Unterschied der loxodromen und der orthodromen Entfernung zwischen dem Nadelkap (Südafrika; $\varphi_1 = 35,0^\circ \text{ S}$, $\lambda_1 = 20,0^\circ \text{ O}$) und dem Südkap auf Tasmanien (Australien; $\varphi_2 = 43,7^\circ \text{ S}$, $\lambda_2 = 146,7^\circ \text{ O}$) rund 1180 km und der gleiche Unterschied zwischen Kap Hoorn (Südamerika; $\varphi_1 = 56,0^\circ \text{ S}$, $\lambda_1 = 67,4^\circ \text{ W}$) und dem Südkap auf der Stewartinsel (Neuseeland; $\varphi_2 = 47,4^\circ \text{ S}$, $\lambda_2 = 167,6^\circ \text{ O}$) rund 1210 km. Liegen zwei Orte auf derselben geographischen Breite, so ist die Loxodrome, die hier

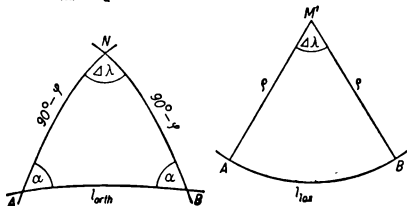
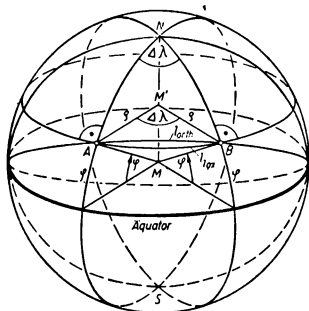


Abb. 27

Wegen $\text{arc } \Delta\lambda = \frac{\pi \Delta\lambda}{180^\circ}$ und $\varrho = r \cdot \cos \varphi$ ergibt sich $l_{\text{lox}} = \frac{\pi \cdot r \cdot \Delta\lambda \cdot \cos \varphi}{180^\circ}$.

Diese spezielle Loxodrome ist eine geschlossene Kurve, und zwar ein Kleinkreis.

Da die loxodrome Entfernung wesentlich größer als die orthodrome Entfernung sein kann, wird im allgemeinen nicht der loxodrome Kurs gefahren. Es werden vielmehr auf dem orthodromen Kurs einzelne Punkte festgelegt. Zwischen diesen Punkten wird dann ein loxodromer Kurs gesteuert. Man braucht dann den Kurs nur in bestimmten Zeitabschnitten und nicht stetig zu ändern. Der Unterschied zwischen der orthodromen Entfernung und einer solchen ist dann nicht mehr erheblich. Wir können die Annäherung der Orthodromen durch Loxodrome auf einer Kugel mit der Annäherung eines Kreises durch einen Polygonzug in der Ebene vergleichen.

Liegen zwei Orte A und B hinreichend nahe beieinander, so ist, wenn α ein spitzer Winkel ist, auch noch $180^\circ - \beta$ ein spitzer, das heißt, β ist ein stumpfer Winkel. Wegen der stetigen Kursänderung auf der Orthodromen wird $180^\circ - \beta$ allmählich immer größer, und nach Erreichen des Wertes

$$\beta = 180^\circ - \beta = 90^\circ \quad \text{wird} \quad 180^\circ - \beta$$

ein stumpfer, das heißt, β wird ein spitzer Winkel.

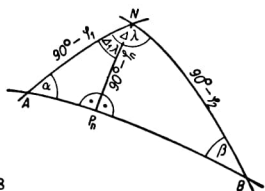
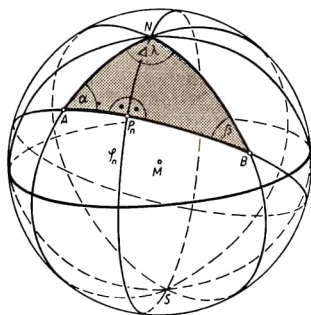


Abb. 28

Für die folgenden Betrachtungen wollen wir ein Poldreieck verwenden, das als dritten Eckpunkt den Nordpol enthält. Sind nun α und β spitze Winkel, so gibt es auf dem Wege AB einen Meridian, der von der Orthodromen unter einem Winkel

von 90° geschnitten wird. Der Schnittpunkt P_n dieses Meridians mit der Orthodromen ist von allen Punkten auf dem Wege AB der dem Nordpol N am nächsten gelegene Punkt. Er heißt **Scheitelpunkt** auf dem Wege AB . Bezeichnen wir seine geographischen Koordinaten mit $P_n(\varphi_n; \lambda_n)$ und setzen $\lambda_n - \lambda_1 = \Delta_1\lambda$, so ist die Seite $P_nN \approx 90^\circ - \varphi_n$ und der Winkel $ANP_n = \Delta_1\lambda$. In dem bei P_n rechtwinkligen sphärischen Dreieck AP_nN kann $\Delta_1\lambda$ und φ_n mit Hilfe der Neperschen Formeln berechnet werden, wenn der Kurswinkel α vorher im sphärischen Dreieck ANB berechnet worden ist (Abb. 28). Ähnliche Überlegungen gelten, wenn α und β stumpfe Winkel sind. Dann ist der Schnittpunkt P_s (der Scheitelpunkt), in dem der Meridian die Orthodrome rechtwinklig schneidet, dem Südpol am nächsten gelegen.

a) Nach dem Seitenkosinussatz ergibt sich die Seite $HV = l$ (Abb. 32 b) aus

$$\cos l = \cos(90^\circ - \varphi_1) \cdot \cos(90^\circ + |\varphi_2|) + \sin(90^\circ - \varphi_1) \cdot \sin(90^\circ + |\varphi_2|) \cdot \cos \Delta_1 \lambda$$

oder

$$\cos l = -\sin \varphi_1 \cdot \sin |\varphi_2| + \cos \varphi_1 \cdot \cos |\varphi_2| \cdot \cos \Delta_1 \lambda. \quad (1)$$

Die Winkel α und β ergeben sich nach dem Sinussatz (Abb. 32 b) aus

$$\sin \alpha = \frac{\sin(90^\circ + |\varphi_2|) \cdot \sin \Delta_1 \lambda}{\sin l}$$

oder

$$\sin \alpha = \frac{\cos |\varphi_2| \cdot \sin \Delta_1 \lambda}{\sin l} \quad (2)$$

und

$$\sin \beta = \frac{\sin(90^\circ - \varphi_1) \cdot \sin \Delta_1 \lambda}{\sin l}$$

oder

$$\sin \beta = \frac{\cos \varphi_1 \cdot \sin \Delta_1 \lambda}{\sin l}. \quad (3)$$

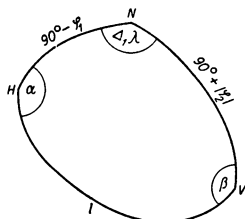


Abb. 32 b

b) Die Berechnung des südlichsten Punktes ergibt sich nach Abbildung 32 c. Das Dreieck NHP ist bei P rechtwinklig. Es ergibt sich zunächst für $\Delta_2 \lambda$ nach der Napierschen Regel

$$\cos(90^\circ - \varphi_1) = \text{ctg} \Delta_2 \lambda \cdot \text{ctg} \alpha,$$

also

$$\text{ctg} \Delta_2 \lambda = \frac{\sin \varphi_1}{\text{ctg} \alpha} = \sin \varphi_1 \cdot \text{tg} \alpha. \quad (4)$$

Dann folgt für die geographische Länge des Punktes P

$$\lambda_s = \lambda_1 + \Delta_2 \lambda. \quad (5)$$

Die geographische Breite des Punktes P ergibt sich aus

$$\cos(-|\varphi_s|) = \sin \alpha \cdot \sin(90^\circ - \varphi_1),$$

$$\cos |\varphi_s| = \sin \alpha \cdot \cos \varphi_1. \quad (6)$$

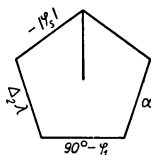
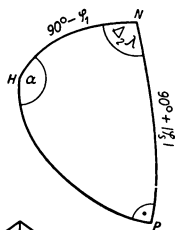


Abb. 32 c

o) Die Berechnung des Schnittpunktes mit der Datumgrenze ergibt sich nach Abb. 32 d. Im Dreieck HND sind die folgenden Seiten und Winkel bekannt: $HN = 90^\circ - \varphi_1$, $\Delta_2 \lambda = 180^\circ - \lambda_1$ und α . Berechnet werden muß die Seite $DN = 90^\circ + |\varphi_D|$. Zunächst ist es notwendig, den Winkel ε zu berechnen.

Es ist nach dem Winkelkosinussatz

$$\cos \varepsilon = -\cos \Delta_3 \lambda \cdot \cos \alpha + \sin \Delta_3 \lambda \cdot \sin \alpha \cdot \cos (90^\circ - \varphi_1)$$

oder

$$\cos \varepsilon = -\cos \Delta_3 \lambda \cdot \cos \alpha + \sin \Delta_3 \lambda \cdot \sin \alpha \cdot \sin \varphi_1. \quad (7)$$

Es ergibt sich nun für $|\varphi_D|$:

$$\sin (90^\circ + |\varphi_D|) = \frac{\sin \alpha \cdot \sin (90^\circ - \varphi_1)}{\sin \varepsilon}$$

oder

$$\cos |\varphi_D| = \frac{\sin \alpha \cdot \cos \varphi_1}{\sin \varepsilon}. \quad (8)$$

Abb. 32 d

- d) Die Berechnung des Schnittpunktes Q mit dem Äquator ergibt sich aus Abbildung 32e. Das Dreieck $QV_A V$ ist bei V_A rechtwinklig. Daraus ergibt sich für $\Delta_4 \lambda$

$$\cos (90^\circ - |\varphi_2|) = \operatorname{ctg} \beta \cdot \operatorname{ctg} (90^\circ - \Delta_4 \lambda)$$

oder

$$\sin |\varphi_2| = \operatorname{ctg} \beta \cdot \operatorname{tg} \Delta_4 \lambda.$$

Es folgt

$$\operatorname{tg} \Delta_4 \lambda = \sin |\varphi_2| \cdot \operatorname{tg} \beta. \quad (9)$$

Dann ist

$$\lambda_A = \lambda_2 - \Delta_4 \lambda. \quad (10)$$

- e) Die Berechnung des Schnittpunktes W mit dem südlichen Wendekreis ergibt sich aus Abbildung 32f. Im Dreieck HNW sind die folgenden Seiten und Winkel bekannt: $HN = 90^\circ - \varphi_1$, $NW = 90^\circ + |\varphi_W|$ und α . Zunächst wird der Winkel γ berechnet. Es ist

$$\sin \gamma = \frac{\sin \alpha \cdot \sin (90^\circ - \varphi_1)}{\sin (90^\circ + |\varphi_W|)}$$

oder

$$\sin \gamma = \frac{\sin \alpha \cdot \cos \varphi_1}{\cos |\varphi_W|}. \quad (11)$$

Abb. 32 f

Zur Berechnung des Winkels $\Delta_5 \lambda$ verwenden wir den Lösungsweg mit Hilfe einer Transformation (vgl. Seite 93). Es ist

$$\cos \gamma = -\cos \alpha \cdot \cos \Delta_5 \lambda + \sin \alpha \cdot \sin \Delta_5 \lambda \cdot \cos (90^\circ - \varphi_1)$$

oder

$$\cos \gamma = -\cos \alpha \cdot \cos \Delta_5 \lambda + \sin \alpha \cdot \sin \Delta_5 \lambda \cdot \sin \varphi_1.$$

¹⁾ Setzt man die Lösung b (Koordinaten des Scheitelpunktes) als bekannt voraus, so läßt sich die Rechnung c im rechtwinkligen Dreieck DPN und die Rechnung e mit Hilfe des Dreiecks WPN wesentlich vereinfachen.

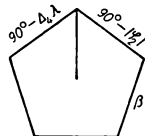
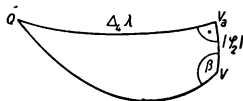
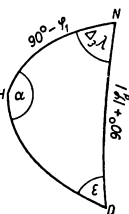
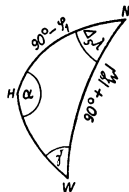


Abb. 32 e



Wir setzen

$$\operatorname{tg} \psi = \operatorname{tg} \alpha \cdot \sin \varphi_1 \quad (12)$$

oder

$$\cos \alpha \cdot \operatorname{tg} \psi = \sin \alpha \cdot \sin \varphi_1.$$

Dann gilt

$$\begin{aligned} \cos \gamma &= -\cos \alpha \cdot \cos \Delta_5 \lambda + \cos \alpha \cdot \frac{\sin \psi}{\cos \psi} \cdot \sin \Delta_5 \lambda, \\ \cos \psi \cdot \cos \gamma &= -\cos \psi \cdot \cos \alpha \cdot \cos \Delta_5 \lambda + \cos \alpha \cdot \sin \psi \cdot \sin \Delta_5 \lambda, \\ &= \cos \alpha \cdot (\sin \psi \cdot \sin \Delta_5 \lambda - \cos \psi \cdot \cos \Delta_5 \lambda), \\ &= -\cos \alpha \cdot \cos (\psi + \Delta_5 \lambda), \end{aligned}$$

also

$$\cos (\psi + \Delta_5 \lambda) = -\frac{\cos \psi \cdot \cos \gamma}{\cos \alpha}. \quad (13)$$

Dann folgt

$$\lambda_W = \lambda_1 + \Delta_5 \lambda. \quad (14)$$

2. Numerische Lösung

a) (1) $\cos l = -\sin \varphi_1 \sin |\varphi_2|$
 $+ \cos \varphi_1 \cos |\varphi_2| \cos \Delta_1 \lambda$

(1a) $\Delta_1 \lambda = 288,3^\circ - 114,2^\circ = 174,1^\circ$

	num	lg
$\sin \varphi_1$	$\sin 22,3^\circ$	9,5792 - 10
$\sin \varphi_2 $	$\sin 33,0^\circ$	9,7361 - 10
	-0,2067	0,3153 - 1

$\cos \varphi_1$	$\cos 22,3^\circ$	9,9662 - 10
$\cos \varphi_2 $	$\cos 33,0^\circ$	9,9236 - 10
$\cos \Delta_1 \lambda$	$\cos 174,1^\circ$	9,9977 - 10

	-0,7718	0,8875 - 1
--	---------	------------

$\cos l$	-0,9785	
----------	---------	--

$$l = 180^\circ - 11,9^\circ = 168,1^\circ$$

$$l = 168,1 \cdot 60 \text{ sm} = 10086 \text{ sm}$$

(2) $\sin \alpha = \frac{\cos |\varphi_2| \sin \Delta_1 \lambda}{\sin l}$

	num	lg
$\cos \varphi_2 $	$\cos 33,0^\circ$	9,9236 - 10
$\sin \Delta_1 \lambda$	$\sin 174,1^\circ$	9,0120 - 10
$\sin l$	$\sin 168,1^\circ$	9,3143 - 10

$\sin \alpha$	$\sin 24,72^\circ$	9,6213 - 10
---------------	--------------------	-------------

$$\alpha = 180^\circ - 24,72^\circ = 155,28^\circ$$

(3) $\sin \beta = \frac{\cos \varphi_1 \sin \Delta_1 \lambda}{\sin l}$

	num	lg
$\cos \varphi_1$	$\cos 22,3^\circ$	9,9662 - 10
$\sin \Delta_1 \lambda$	$\sin 174,1^\circ$	9,0120 - 10
$\sin l$	$\sin 168,1^\circ$	9,3143 - 10
$\sin \beta$	$\sin 27,47^\circ$	9,6639 - 10

$$\beta = 180^\circ - 27,47^\circ = 152,53^\circ$$

Die Winkel α und β sind stumpfe Winkel, wie man leicht am Globus nachprüfen kann.

b) (4) $\operatorname{ctg} \Delta_2 \lambda = \sin \varphi_1 \operatorname{tg} \alpha$

	num	lg
$\sin \varphi_1$	$\sin 22,3^\circ$	9,5792 - 10
$\operatorname{tg} \alpha$	$\operatorname{tg} 155,28^\circ$	9,6631 - 10
$\operatorname{ctg} \Delta_2 \lambda$	$-\operatorname{ctg} 80,09^\circ$	9,2423 - 10

$$\Delta_2 \lambda = 180^\circ - 80,09^\circ = 99,91^\circ$$

(5) $\lambda_s = \lambda_1 + \Delta_2 \lambda$

$$\lambda_s = 114,2^\circ + 99,91^\circ = 214,11^\circ \text{ O}$$

$$\lambda_s = 145,89^\circ \text{ W}$$

$$(6) \quad \cos |\varphi_S| = \sin \alpha \cos \varphi_1$$

	num	lg
$\sin \alpha$	$\sin 155,28^\circ$	9,6213 - 10
$\cos \varphi_1$	$\cos 22,3^\circ$	9,9662 - 10
$\cos \varphi_S $	$\cos 67,25^\circ$	9,5875 - 10

$$\varphi_S = 67,25^\circ \text{ S}$$

Die Koordinaten des südlichsten Punktes sind $P (67,25^\circ \text{ S}; 145,89^\circ \text{ W})$.

$$o) \quad (7) \quad \cos \varepsilon = -\cos \Delta_3 \lambda \cos \alpha + \sin \Delta_3 \lambda \sin \alpha \sin \varphi_1$$

$$(7a) \quad \Delta_3 \lambda = 180^\circ - 114,2^\circ = 65,8^\circ$$

	num	lg
$\cos \Delta_3 \lambda$	$\cos 65,8^\circ$	9,6127 - 10
$\cos \alpha$	$\cos 155,28^\circ$	9,9582 - 10
	+ 0,3723	0,5709 - 1
$\sin \Delta_3 \lambda$	$\sin 65,8^\circ$	9,9601 - 10
$\sin \alpha$	$\sin 155,28^\circ$	9,6213 - 10
$\sin \varphi_1$	$\sin 22,3^\circ$	9,5792 - 10
	0,1447	0,1606 - 1
$\cos \varepsilon$	0,5170	

$$\varepsilon = 58,87^\circ$$

$$(8) \quad \cos |\varphi_D| = \frac{\sin \alpha \cos \varphi_1}{\sin \varepsilon}$$

	num	lg
$\sin \alpha$	$\sin 155,28^\circ$	9,6213 - 10
$\cos \varphi_1$	$\cos 22,3^\circ$	9,9662 - 10
$\sin \varepsilon$	$\sin 58,87^\circ$	9,9325 - 10
$\cos \varphi_D $	$\cos 63,14^\circ$	9,6550 - 10

$$\varphi_D = 63,14^\circ \text{ S}$$

Die Datumgrenze wird unter der Breite $63,14^\circ \text{ S}$ geschnitten.

$$d) \quad (9) \quad \lg \Delta_4 \lambda = \lg \sin |\varphi_2| \lg \beta$$

	num	lg
$\sin \varphi_2 $	$\sin 33^\circ$	9,7361 - 10
$\lg \beta$	$\lg 152,53^\circ$	9,7159 - 10
$\lg \Delta_4 \lambda$	$-\lg 15,81^\circ$	9,4520 - 10
	$\Delta_4 \lambda = 164,19^\circ$	

$$(10) \quad \lambda_A = \lambda_2 - \Delta_4 \lambda = 283,3^\circ - 164,19^\circ = 124,11^\circ$$

Der Äquator wird unter der Länge $124,11^\circ \text{ O}$ geschnitten.

$$e) \quad (11) \quad \sin \gamma = \frac{\sin \alpha \cos \varphi_1}{\cos |\varphi_W|}$$

	num	lg
$\sin \alpha$	$\sin 155,28^\circ$	9,6213 - 10
$\cos \varphi_1 $	$\cos 22,3^\circ$	9,9662 - 10
$\cos \varphi_W $	$\cos 23,5^\circ$	9,9624 - 10
$\sin \gamma$	$\sin 24,95^\circ$	9,6251 - 10
	$\gamma = 24,95^\circ$	

$$(12) \quad \lg \psi = \lg \alpha \sin \varphi_1$$

	num	lg
$\lg \alpha$	$\lg 155,28^\circ$	9,6631 - 10
$\sin \varphi_1$	$\sin 22,3^\circ$	9,5792 - 10
$\lg \psi$	$-\lg 9,91^\circ$	9,2423 - 10

$$\psi = 180^\circ - 9,91^\circ = 170,09^\circ$$

$$(13) \quad \cos (\psi + \Delta_5 \lambda) = -\frac{\cos \psi \cos \gamma}{\cos \alpha}$$

	num	lg
$\cos \psi$	$\cos 170,09^\circ$	9,9935 - 10
$\cos \gamma$	$\cos 24,95^\circ$	9,9574 - 10
$\cos \alpha$	$\cos 155,28^\circ$	9,9582 - 10
$\cos (\psi + \Delta_5 \lambda)$	$-\cos 10,5^\circ$	9,9927 - 10

$$\varphi + \Delta_\epsilon \lambda = 190,5^\circ$$

$$\Delta_\epsilon \lambda = 190,5^\circ - 170,09^\circ = 20,41^\circ$$

Die Rechnung mit dem Wert $\varphi + \Delta_\epsilon \lambda = 169,5^\circ$ würde einen negativen Wert für $\Delta_\epsilon \lambda$ ergeben.

$$(14) \lambda_W = \lambda_1 + \Delta_\epsilon \lambda = 114,2^\circ + 20,41^\circ$$

$$\lambda_W = 134,61^\circ$$

Der südliche Wendekreis wird unter der Länge $134,61^\circ$ O geschnitten.

Aufgaben

Geographische Koordinaten einiger Orte:

Ort	Geogr. Breite (φ)	Geogr. Länge (λ)	Ort	Geogr. Breite (φ)	Geogr. Länge (λ)
Athen	38,0° N	23,7° O	Neapel . .	40,8° N	14,3° O
Auckland (Neusee- land)	36,8° S	174,8° O	New York	40,8° N	74,0° W
Berlin .	52,5° N	13,4° O	Oslo .	59,9° N	10,7° O
Bombay .	18,9° N	72,8° O	Paris	48,8° N	2,3° O
Budapest	47,5° N	19,1° O	Peking .	39,9° N	116,4° O
Charkow .	50,0° N	36,2° O	Potsdam . . .	52,4° N	13,1° O
Greenwich	51,5° N	0,0°	Rio de Janeiro	22,9° S	43,2° W
Hamburg	53,6° N	10,0° O	Rostock . .	54,1° N	12,1° O
Hongkong	22,3° N	114,2° O	San Francisco	37,8° N	122,4° W
Kairo . .	30,1° N	31,3° O	Strasbourg . . .	48,6° N	7,8° O
Kap Hoorn . .	56,0° S	67,3° W	Südkap (Stewart- Island)	47,3° S	167,5° O
Kap Oljutorski	59,9° N	170,5° O	Südwestkap		
Kapstadt	33,9° S	18,5° O	(Tasmanien)	43,6° S	146,1° O
Leipzig .	51,3° N	12,4° O	Sydney . .	33,9° S	151,2° O
Leningrad	59,9° N	30,3° O	Taschkent	41,3° N	69,2° O
London .	51,5° N	0,2° W	Tbilissi .	41,7° N	44,8° O
Melbourne .	38,5° S	144,7° O	Tokio . .	35,7° N	139,8° O
Moskau . .	55,8° N	37,6° O	Ulan-Bator .	47,8° N	106,9° O
Nadelkap	34,8° S	20,0° O	Valparaiso .	33,0° S	71,7° W
Nagasaki	32,8° N	129,9° O	Wladiwostok . . .	42,9° N	132,0° O

In den folgenden Aufgaben sind die geographischen Koordinaten, wenn nicht besonders angegeben, dieser Tabelle zu entnehmen.

- Wie groß ist die kürzeste Entfernung zwischen Charkow und Peking, und wie groß sind die Kurswinkel?
- Das im Jahre 1874 von der Insel Valentia ($\varphi_1 = 51,9^\circ$ N; $\lambda_1 = 10,4^\circ$ W) nach Neufundland ($\varphi_2 = 47,7^\circ$ N; $\lambda_2 = 53,4^\circ$ W) verlegte Kabel hat eine Länge von 1854 sm. Vergleichen Sie diese Länge mit der kürzesten Entfernung zwischen den beiden Orten!

8. Es ist die Länge des Weges zwischen dem Nadelkap und dem Südwestkap von Tasmanien zu bestimmen. Der orthodrome Kurs schneidet den 55. südlichen Breitenparallel in den Punkten E_1 und E_2 . Da ein Schiff die Packeiszone nicht durchfahren darf, fährt es auf dem orthodromen Kurs bis E_1 , dann auf dem Breitenparallel bis E_2 und dann wieder auf dem orthodromen Kurs. Die Kurswinkel sind zu bestimmen.
4. Wo liegt der Scheitelpunkt auf dem kürzesten Wege von Rostock nach Wladiwostok und wie weit ist er vom Nordpol entfernt?
5. Wie groß ist der Unterschied zwischen der loxodromen und der orthodromen Entfernung
 - a) Leningrad—Oslo,
 - b) New York—Neapel,
 - c) Kapstadt—Sydney,
 - d) Leningrad—Kap Oljutorski?

Wie weit ist der nördlichste bzw. der südlichste Punkt von der Loxodromen entfernt? (Die Loxodrome liegt in diesen Fällen auf dem entsprechenden Breitenkreis.)
6. Es sind die kürzeste Entfernung für die Fluglinie zwischen Moskau und Peking über Ulan-Bator und die Kurswinkel zu berechnen. Welches ist der nördlichste Punkt jeder Teilstrecke?
7. Der Portugiese Magalhães benötigte bei seiner Erdumseglung 1520/21 für die Entfernung von der nach ihm benannten Meeresstraße ($\varphi_1 = 54,5^\circ \text{ S}$; $\lambda_1 = 71,5^\circ \text{ W}$) nach den Philippinen ($\varphi_2 = 8,0^\circ \text{ N}$; $\lambda_2 = 126,3^\circ \text{ O}$) 99 Tage. Welche Zeit brauchte ein Schiff, das mit einer Geschwindigkeit von 18 Knoten (1 Knoten = 1 sm/h) fährt, wenn es zunächst zum Ort P ($\varphi_3 = 8,6^\circ \text{ N}$; $\lambda_3 = 170,0^\circ \text{ O}$) auf der Orthodromen und dann weiter auf der Loxodromen fahren würde?
8. Ein Schiff fährt von Auckland (Neuseeland) mit dem orthodromen Kurs *ONO* ab. Wo und wann kreuzt es die Datumgrenze und wo und wann den Äquator, wenn das Schiff mit einer Geschwindigkeit von 19,4 Knoten (1 Knoten = 1 sm/h) fährt?
9. Es sind die kürzeste Entfernung zwischen Leningrad und Wladiwostok und der Scheitelpunkt zu bestimmen.
10. Von einem Ort P ($\varphi_1 = 34^\circ 20' \text{ S}$; $\lambda_1 = 18^\circ 20' \text{ O}$) will man auf dem kürzesten Wege nach Melbourne fahren, ohne wegen der Eisgefahr den Breitenparallel 55° S zu überschreiten.
 - a) Auf welchen Längen wird dieser Breitenparallel erreicht und verlassen?
 - b) Wieviel Seemeilen sind im ganzen zurückzulegen?
 - c) In welchen Breiten schneiden die beiden Großkreisbögen die durch 5 teilbaren Meridiane?
 - d) Bestimmen Sie den Anfangskurs und den Endkurs!
11. Ein Schiff befindet sich in einem Ort P ($\varphi_1 = 50^\circ 10' \text{ S}$; $\lambda_1 = 159^\circ 20' \text{ W}$) und soll, ohne in südlichere Breiten zu kommen, auf dem kürzesten Weg nach Valparaiso fahren.
 - a) Auf welcher Länge ist der Breitenparallel $50^\circ 10' \text{ S}$ zu verlassen?
 - b) Wieviel Seemeilen sind im ganzen zurückzulegen?
 - c) Mit welchem Endkurs kommt man in Valparaiso an?
 - d) Die Schnittpunkte der Orthodromen mit den durch 5 teilbaren Meridianen sind zu berechnen.

12. Es sind die Kurswinkel und die kürzeste Entfernung für eine Fluglinie zu berechnen.
- | | |
|---------------------|-----------------------|
| a) Leningrad—Moskau | b) Berlin—Budapest |
| c) Moskau—Tbilissi | d) Peking—Wladiwostok |
13. Berechnen Sie die Orthodrome zwischen Kapstadt und Rio de Janeiro und geben Sie die Kurswinkel und den Scheitelpunkt der Orthodromen an!
14. Von einem Ort P ($\lambda_1 = 136,9^\circ \text{ W}$; $\varphi_1 = 35,8^\circ \text{ N}$) soll ein Schiff auf dem kürzesten Wege nach San Francisco fahren. Es sind die kürzeste Entfernung, der Anfangskurs, der Endkurs und der Scheitelpunkt zu berechnen.
15. Ein Schiff soll von Kapstadt nach Bombay fahren. Es fährt zunächst auf dem Längengrad bis zum Breitenparallel 35° S und dann auf dem Breitenparallel bis 48° östlicher Länge. Von dort aus fährt es auf der Orthodrome bis Bombay. Wie lang ist der Weg? Welches sind die Kurswinkel in den einzelnen Orten? Um wieviel ist dieser Weg länger als die kürzeste Entfernung zwischen Kapstadt und Bombay?
16. Welche beiden Punkte des Äquators sind von Taschkent ebenso weit entfernt wie Taschkent vom Nordpol?
17. Es ist die kürzeste Entfernung Kairo—Leipzig zu bestimmen.
18. Ein Schiff soll von Valparaiso nach Südkap (Stewart-Insel*) fahren, ohne den Breitenparallel 55° S zu überschreiten. Wie lang ist der Weg? Welche Kurswinkel ergeben sich? Um wieviel ist der Weg länger als die kürzeste Entfernung? Wie weit ist der Scheitelpunkt der Orthodrome vom Südpol entfernt?
19. Es ist der kürzeste Weg zwischen Taschkent und Leningrad zu bestimmen. Unter welchen Längen (Breiten) werden die durch 5 teilbaren Breitenkreise (Längengrade) geschnitten? Kennzeichnen Sie den Kurs auf der Karte!
20. Auf einer Karte ist die Orthodrome zwischen Charkow und Kap Oljutorski einzutragen.
21. Auf einer Mercatorkarte (Atlas zur Erd- und Länderkunde, 92c—e) sind die Loxodrome und die Orthodrome zwischen dem Südwestkap (Tasmanien) und Neapel einzutragen.

Anmerkung zu den Aufgaben 20 und 21: Es sind mehrere Punkte der Orthodromen zu berechnen.

D. Komplexe Zahlen

I. Vom Bereich der natürlichen Zahlen zum Bereich der reellen Zahlen

Der heutige Zahlbegriff ist das Ergebnis einer langen Entwicklung, die bereits in vorgeschichtlicher Zeit begann. Aus den praktischen Bedürfnissen der Menschen ergaben sich das Zählen und damit die ersten Zahlen der Folge

1, 2, 3, 4, 5, ...

Sie sind Zahlen im ursprünglichen Sinne des Wortes, das heißt, man kann mit ihnen zählen. Doch kennzeichnet das Zählen nur einen Teil ihrer Bedeutung. Es steht nämlich in einem engen wechselseitigen Zusammenhang mit dem Rechnen. So entwickeln sich einerseits die ersten Rechenoperationen unmittelbar aus dem Zählen durch Zusammenfassung einzelner Zählsschritte, während andererseits das Zählen ohne die Hilfe des Rechnens nach wenigen Schritten aufhören würde. Wir können zu größeren Zahlen nur gelangen, indem wir kleinere rechnend zusammensetzen. Bei allen Zahlen von 11 an kommt dies bereits in unserer Ziffernschreibweise zum Ausdruck (vgl. Aufgabe 1). Erst das Rechnen ermöglicht es also, die Zahlen zu beherrschen und sie im praktischen Leben anzuwenden. Die wesentliche Bedeutung der Zahlen liegt also nicht darin, daß man mit ihnen zählen, sondern daß man mit ihnen rechnen kann.

Um so entscheidender ist die Tatsache, daß mit den betrachteten Zahlen wohl das Zählen ohne jede Einschränkung möglich ist, jedoch viele von den einfachsten Rechenaufgaben mit ihnen nicht lösbar sind. Es ist daher erforderlich, umfassendere Zahlbereiche zu bilden, in denen solche bisher unlösbaren Aufgaben gelöst werden können. Auch bei diesen Erweiterungen ist es berechtigt, von Zahlen zu sprechen; denn schon bei den ursprünglichen Zahlen 1, 2, 3, ... ist das Rechnen wesentlicher als das Zählen. Wir müssen nur zuvor den ursprünglichen Zahlen einen besonderen Namen geben, damit wir sie von dem erweiterten Begriff unterscheiden können. Die bekannte Bezeichnung „natürliche Zahlen“ soll daran erinnern, daß diese Zahlen die erste Stufe des Zahlbegriffes darstellen, daß es Zahlen sind, mit denen man nicht nur rechnen, sondern auch zählen kann.

Von den Erweiterungen des Zahlbegriffes sind uns schon einige bekannt. In der Grundschule haben wir ganze und gebrochene, positive und negative Zahlen kennen gelernt. Im 9. Schuljahr wurden als Erweiterung der rationalen Zahlen die reellen Zahlen behandelt. Wie wir im folgenden feststellen werden, läßt sich der Bereich der reellen Zahlen noch zu dem der komplexen Zahlen erweitern.

Aufgabe

Jede Ziffernschreibweise natürlicher Zahlen beruht – von den grundlegenden Zeichen für die ersten Zahlen abgesehen – auf Rechenoperationen.

- Welche Rechenschritte führen zum Beispiel zu der mit dem Symbol 2 645 gekennzeichneten natürlichen Zahl?
- Welche Rechenoperationen liegen also der dekadischen Schreibweise zugrunde?
- Welche Rechenoperationen sind dagegen für die Schreibweise im römischen Ziffernsystem notwendig, etwa für die Zahl CCXLIV?
- Welche Rechenoperationen erfordert schließlich die Schreibweise einer gebrochenen Zahl, etwa 17,125 im dekadischen System?

1. Der Bereich der natürlichen Zahlen

Durch das Zählen erhalten die natürlichen Zahlen von selbst eine Ordnung, die wir ihre lineare Anordnung nennen wollen. Die natürlichen Zahlen sind in einer ganz bestimmten Reihenfolge gegeben, in der jede von ihnen genau einen unmittelbaren Nachfolger und genau einen unmittelbaren Vorgänger hat. Davon ausgenommen ist nur die 1, die nur Nachfolger, aber keine Vorgänger hat.

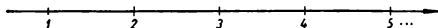


Abb. 33

Diese Eigenschaft der natürlichen Zahlen ermöglicht die Veranschaulichung auf dem Zahlenstrahl (Abb. 33). Dabei wird eine Zahl selbst auch oft als der Weg (Pfeil) von einem Anfangspunkt A zu dem mit ihr bezeichneten Punkt angesehen (Abb. 34).

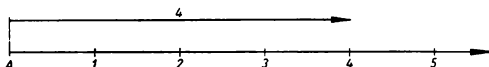


Abb. 34

Der Abstand von Punkt zu Punkt bedeutet dann immer das Weitergehen von der einen Zahl zur nächstfolgenden, es gibt also einen Schritt an.

a) Die Addition natürlicher Zahlen

Die Zusammenfassung mehrerer Zählsschritte führt zur ersten Rechenoperation, der Addition. Die Summe zweier natürlicher Zahlen a und b ist die Zahl, die man erhält, wenn man von a aus b Schritte weiterzählt. Für die natürlichen Zahlen ergeben sich die folgenden

Grundgesetze der Addition:

1. Die Summe zweier natürlicher Zahlen ist stets wieder eine eindeutig bestimmte natürliche Zahl. Die Addition ist also im Bereich der natürlichen Zahlen immer ausführbar.
2. Die Addition ist kommutativ, das heißt, es gilt

$$a + b = b + a.$$

3. Die Addition ist assoziativ, das heißt, es gilt

$$a + (b + c) = (a + b) + c.$$

Diese Grundgesetze sind aus der Rechenpraxis von Jahrtausenden abstrahiert, so daß sich zunächst jeder Zweifel an ihnen verbietet. Man kann aber einwenden, daß es sicher sehr große natürliche Zahlen gibt, mit denen bisher noch niemand gerechnet hat. Es ist deshalb von Bedeutung, daß man die uneingeschränkte Gültigkeit dieser Grundgesetze aus der Erklärung der Addition mit Hilfe der linearen Anordnung ableiten kann (vgl. Aufgabe 1). Die Berechtigung, gerade diese drei Regeln als Grundgesetze der Addition auszuzeichnen, liegt darin, daß sich alle übrigen Rechenregeln der Addition als Folgerungen aus ihnen ableiten lassen (vgl. Aufgabe 2).

b) Die Multiplikation natürlicher Zahlen

Wie die Addition aus dem Zusammenfassen des Zählens entsteht, so ergibt sich die Multiplikation aus dem Zusammenfassen gewisser Additionsaufgaben. Es bedeutet $a \cdot b$ nichts anderes als $b + b + \dots + b$, wobei b genau a mal als Summand gesetzt wird. Diese Tatsache rechtfertigt auch die Bezeichnung Rechenoperation zweiter Stufe. Es gelten dabei analog die folgenden

Grundgesetze der Multiplikation:

1. Das Produkt zweier natürlicher Zahlen ist stets wieder eine eindeutig bestimmte natürliche Zahl. Die Multiplikation ist also im Bereich der natürlichen Zahlen immer ausführbar.
2. Die Multiplikation ist kommutativ, das heißt, es gilt

$$a \cdot b = b \cdot a.$$

3. Die Multiplikation ist assoziativ, das heißt, es gilt

$$a \cdot (b \cdot c) = (a \cdot b) \cdot c.$$

Dabei stehen beide Rechenoperationen in einem Zusammenhang, der zum Ausdruck kommt in dem

Grundgesetz der Verknüpfung beider Operationen:

Die Addition und die Multiplikation sind distributiv verknüpft, das heißt, es gilt

$$a \cdot (b + c) = a \cdot b + a \cdot c.$$

Für diese vier letztgenannten Grundgesetze gilt das Entsprechende wie für die Grundgesetze der Addition. Auch sie lassen sich aus der linearen Anordnung der natürlichen Zahlen, der Erklärung der Addition und der Erklärung der Multiplikation ableiten (vgl. Aufgabe 3 und 4).

Das gesamte Rechnen im Bereich der natürlichen Zahlen ist durch diese sieben Grundgesetze bereits vollständig bestimmt. Alle übrigen Rechengesetze kann man als Folgerungen aus ihnen gewinnen. Das gilt nicht nur für alle Regeln der Verknüpfung von Addition und Multiplikation (vgl. Aufgabe 5), sondern auch für die Rechenoperation dritter Stufe, das Potenzieren (vgl. Aufgabe 6). Es gilt darüber hinaus sogar für die Rechenoperationen, die aus den Umkehrungen der bisher besprochenen Operationen entstehen, soweit diese überhaupt innerhalb der natürlichen Zahlen eindeutig ausführbar sind.

c) Umkehrung der Addition

Die Addition zweier natürlicher Zahlen ergibt wieder eine eindeutig bestimmte natürliche Zahl. Die Umkehrung besteht in der Aufgabe, zu einer gegebenen natürlichen Zahl b eine zweite Zahl x zu finden, so daß die Summe dieser beiden Zahlen gleich einer gegebenen natürlichen Zahl a ist. Wir müssen also die Gleichung

$$b + x = a$$

lösen, und das ist innerhalb der natürlichen Zahlen nur dann möglich, wenn a größer als b ist. In diesem Fall gibt es genau eine natürliche Zahl x , die der Gleichung genügt und die wir als Differenz

$$x = a - b$$

bezeichnen. In allen anderen Fällen ist die Subtraktion innerhalb der natürlichen Zahlen unlösbar.

d) Umkehrung der Multiplikation

Hier verläuft die Überlegung analog. Auch die Multiplikation zweier natürlicher Zahlen hat als Ergebnis eine eindeutig bestimmte natürliche Zahl. Die Umkehrung besteht in der Aufgabe, zu einer gegebenen natürlichen Zahl b eine zweite Zahl x zu finden, so daß das Produkt dieser beiden Zahlen gleich einer gegebenen natürlichen Zahl a ist. Wir müssen also die Gleichung

$$b \cdot x = a$$

lösen, und das ist innerhalb der natürlichen Zahlen nur dann möglich, wenn a ein Vielfaches von b oder gleich b ist. Dann gibt es genau eine natürliche Zahl x , die diese Gleichung erfüllt und die wir den Quotienten

$$x = a : b = \frac{a}{b}$$

nennen. In allen anderen Fällen ist die Division innerhalb der natürlichen Zahlen unlösbar.

e) Umkehrung des Potenzierens

Die Potenz einer natürlichen Zahl mit einer natürlichen Zahl als Exponenten ist ebenfalls eine eindeutig bestimmte natürliche Zahl. Hier ergeben sich jedoch für die Umkehrung zwei Möglichkeiten (vgl. Aufgabe 7). Man kann einmal nach einer Zahl x fragen, von der eine durch n bestimmte Potenz eine vorgegebene natürliche Zahl a ist, zum anderen danach, mit welchem Exponenten x eine gegebene natürliche Zahl a die Potenz der natürlichen Zahl b ist. Die entsprechenden Gleichungen

$$x^n = a \quad \text{bzw.} \quad b^x = a$$

sind innerhalb der natürlichen Zahlen nur bei geeigneter Wahl von n und a bzw. b und a lösbar. Man schreibt für diese Lösungen, die dann eindeutig sind,

$$x = \sqrt[n]{a} \quad \text{bzw.} \quad x = {}^b\log a$$

und nennt die zugehörigen Operationen Radizieren und Logarithmieren (vgl. Aufgabe 8).

Damit ist der Weg zu allen Erweiterungen des Zahlbegriffes bereits durch die natürlichen Zahlen vorgezeichnet. Um die Subtraktion uneingeschränkt ausführen zu können, muß man die negativen Zahlen einführen. Für die uneingeschränkte Ausführbarkeit der Division ist die Einführung der Brüche erforderlich. Das Radizieren und Logarithmieren führt schließlich zu den irrationalen und komplexen Zahlen.

Aufgaben

1. Die im Text erwähnte Ableitung der drei Grundgesetze der Addition aus der Erklärung der Addition mit Hilfe der linearen Anordnung ist, exakt durchgeführt, verhältnismäßig schwierig. Man kann diese Ableitung aber unter Zuhilfenahme der auf Seite 110 erwähnten Pfeildarstellung veranschaulichen. Überzeugen Sie sich auf diese Weise von der Richtigkeit dieser drei Grundgesetze!
2. Mit Hilfe der drei Grundgesetze der Addition ist zu zeigen, daß es bei der Berechnung einer Summe mit endlich vielen Summanden $a + b + \dots + k$ weder auf die Reihenfolge noch auf eventuell vorgeschriebene Zusammenfassungen (Klammern) ankommt¹⁾.
3. Die drei Grundgesetze der Multiplikation von natürlichen Zahlen sind entsprechend den Überlegungen in Aufgabe 1 mit Hilfe geeigneter Pfeildarstellungen zu bestätigen.

¹⁾ Für eine beliebige endliche Anzahl von Summanden ist ein Induktionsschluß notwendig.

114 Vom Bereich der natürlichen Zahlen zum Bereich der reellen Zahlen

Anleitung: Für das kommutative Gesetz ist eine Zerlegung des Gesamtpfeils (Abb. 35a) in ein rechteckiges Schema zu wählen (Abb. 35b). Für das assoziative Gesetz ist ein entsprechendes dreidimensionales, perspektivisches Schema zu zeichnen.

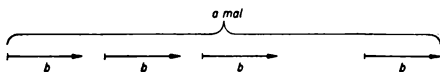


Abb. 35a

4. Entsprechend der Aufgabe 3 ist die Richtigkeit des distributiven Gesetzes für natürliche Zahlen nachzuweisen.

5. Allein mit Hilfe der sieben Grundgesetze sind die folgenden Rechenregeln zu beweisen:

a) $(b + c) \cdot a = ba + ca$

b) $a \cdot (b + c + d) = ab + ac + ad$

c) $(a + b) \cdot (c + d) = ac + ad + bc + bd$

d) Geben Sie die weitestgehende Verallgemeinerung des distributiven Gesetzes an! Wie wäre dieses zu beweisen?

6. Aus der Erklärung des Potenzierens

$a^n = a \cdot a \cdot \dots \cdot a$ (n Faktoren) und mit den sieben Grundgesetzen sind die folgenden Regeln zu beweisen:

a) $a^n \cdot a^m = a^{n+m}$,

b) $a^n \cdot b^n = (ab)^n$,

c) $(a^n)^m = a^{n \cdot m}$.

7. Begründen Sie, weshalb es beim Potenzieren zwei Möglichkeiten für die Umkehrung gibt und weshalb dies bei der Addition und Multiplikation nicht der Fall ist!

8. Lösen Sie die Gleichungen $x^3 = 8$ bzw. $5^x = 125$ und schreiben Sie die Lösungen in der im Text angegebenen Form! Bilden Sie weitere Beispiele dieser Art!

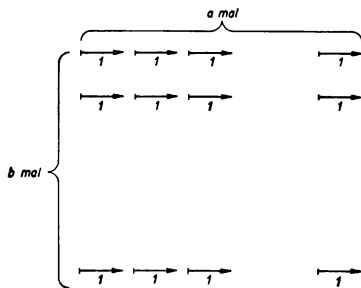


Abb. 35b

2. Die Erweiterung von Zahlenbereichen

Wir wollen zunächst untersuchen, wie ein bereits bekannter Zahlenbereich erweitert werden kann und was man dabei unter „erweitern“ versteht. Als Beispiel wählen wir dazu den Übergang von den natürlichen Zahlen zu den Brüchen, deren Zähler und Nenner natürliche Zahlen sind.

Die Einführung der Brüche und der Bruchrechnung entspringt dem Bedürfnis, jede aus natürlichen Zahlen gebildete Divisionsaufgabe lösen zu können, so zum

Beispiel die Aufgabe $2:3$. Wir erhalten diese innerhalb der natürlichen Zahlen nicht vorhandene Lösung, indem wir

$$2:3 = \frac{2}{3}$$

schreiben. Auf diese Weise erklären wir den Bruch „zwei Drittel“. Wenn wir damit auch eine sehr anschauliche Vorstellung verbinden, so haben wir doch eigentlich nur die innerhalb der natürlichen Zahlen unlösbare Aufgabe nochmals in anderer Gestalt hingeschrieben.

Daß wir damit überhaupt etwas erreicht haben, ist in drei Tatsachen begründet:

1. Mit der Gesamtheit der so erklärten Brüche kann man nach den bekannten Regeln der Bruchrechnung wie mit natürlichen Zahlen rechnen. Es gelten nämlich für die Addition und Multiplikation von Brüchen die gleichen sieben Grundgesetze wie für natürliche Zahlen (Aufgabe 2). Diese Grundgesetze bestimmen aber bereits alle Rechengesetze. Wenn also die gleichen Grundgesetze gelten, so stimmen auch alle daraus abgeleiteten Gesetze überein.
2. Diese Brüche enthalten die natürlichen Zahlen als Brüche mit dem Nenner 1. Die Bruchrechnung umfaßt also das Rechnen mit natürlichen Zahlen.
3. Innerhalb der Brüche ist nicht nur jede Divisionsaufgabe mit natürlichen Zahlen, sondern sogar jede Divisionsaufgabe mit Brüchen ausführbar. Dieser Sachverhalt wird durch ein weiteres Grundgesetz gekennzeichnet: Innerhalb der Brüche mit natürlichen Zahlen als Zähler und Nenner ist die Multiplikation immer umkehrbar, das heißt, jede Divisionsaufgabe ist lösbar.

Zusammenfassend können wir also folgendes feststellen: Mit Hilfe der schon bekannten natürlichen Zahlen werden durch die Bildung der Brüche neue Zahlen abgeleitet. Dies geschieht, indem die Divisionsaufgaben mit natürlichen Zahlen in geeigneter Gestalt, eben als Brüche, geschrieben werden. Das Rechnen mit diesen neuen Zahlen wird durch die Regeln der Bruchrechnung auf das Rechnen mit natürlichen Zahlen zurückgeführt (vgl. Aufgabe 1). Dabei bleiben alle Rechengesetze der natürlichen Zahlen für die neuen Zahlen gültig. Der Bereich der neuen Zahlen ist eine Erweiterung des Ausgangsbereiches der natürlichen Zahlen, da er diese enthält. Der Bereich der neuen Zahlen ist von dem Mangel des Ausgangsbereiches, daß die Division nur beschränkt ausführbar ist, frei.

Mit der Verallgemeinerung dieses Verfahrens haben wir bereits das Prinzip jeder Erweiterung eines vorliegenden Zahlenbereiches gefunden: Gerade die im alten Bereich unlösbaren Aufgaben werden zur Erklärung der neuen Zahlen verwendet. Das Rechnen mit ihnen wird so auf das bereits bekannte Rechnen mit den alten Zahlen zurückgeführt, daß dabei alle Rechengesetze des alten Bereiches gültig bleiben. Der neue Bereich muß den alten enthalten und von dem entsprechenden Mangel des alten Bereiches frei sein.

Die Festlegung, daß die Rechengesetze des alten Bereiches auch im neuen Bereich gelten sollen, nennt man das **Permanenzprinzip** oder das **Prinzip von der Permanenz der Rechengesetze**.

Wir werden in den folgenden Abschnitten erkennen, daß alle Erweiterungen von Zahlenbereichen nach dem eben geschilderten Verfahren durchgeführt werden.

Allerdings ist eine exakte Durchführung dieser Gedankengänge oft zu umfangreich und schwierig, als daß wir sie vollständig behandeln könnten. Wir wollen an Hand unseres Beispiels nur zwei Probleme betrachten, die bei allen Erweiterungen eine Rolle spielen:

Jede der neuen Zahlen kann in unendlich vielen verschiedenen Formen auftreten. So erklären zum Beispiel die Divisionsaufgaben

$$\frac{2}{3}; \quad \frac{4}{6}; \quad \frac{6}{9}; \quad \frac{10}{15}; \quad \frac{14}{21} \quad \text{usw.}$$

den gleichen Bruch. Man muß also kennzeichnen, welche der neuen Zahlen einander gleich sind. Bei den Brüchen geschieht dies durch die Regeln des Kürzens und Erweiterns. Die Rechenregeln der neuen Zahlen müssen selbstverständlich von dieser Mehrdeutigkeit der Schreibweise unabhängig sein (vgl. Aufgabe 3). So darf es keinen Einfluß haben, wenn man beispielsweise $\frac{4}{6}$ statt $\frac{2}{3}$ bei Rechnungen verwendet.

Weiterhin sind die alten Zahlen in den neuen zunächst in anderer Gestalt enthalten. So besteht doch ein begrifflicher Unterschied zwischen dem Bruch $\frac{5}{1}$ und der natürlichen Zahl 5, von dem wir gerade absehen, wenn wir sagen, daß die Brüche die natürlichen Zahlen enthalten. Dies ist möglich, da das Rechnen mit allen Brüchen $\frac{a}{1}, \frac{b}{1}, \frac{c}{1}, \dots$ völlig gleich verläuft zu dem Rechnen mit den natürlichen Zahlen $a, b, c, \dots$. In der höheren Mathematik sagt man für diesen Sachverhalt: Die Brüche mit dem Nenner 1 sind ein isomorphes Bild der natürlichen Zahlen und können durch diese ersetzt werden. Erst nach dieser Ersetzung wird aus dem völlig neuen Rechenbereich der Brüche eine Erweiterung des Bereiches der natürlichen Zahlen.

Den Bereich der gebrochenen Zahlen kann man auf dem Zahlenstrahl veranschaulichen, indem man die Einheitsschritte entsprechend unterteilt. Es zeigt sich dann, daß gleichen Brüchen (z. B. $\frac{2}{3}$ und $\frac{4}{6}$) gleiche Punkte entsprechen.

Aufgaben

1. Überlegen Sie, daß die Addition und die Multiplikation von Brüchen durch die Rechenregeln

$$\frac{a}{b} + \frac{c}{d} = \frac{a \cdot d + b \cdot c}{b \cdot d} \quad \text{und} \quad \frac{a}{b} \cdot \frac{c}{d} = \frac{a \cdot c}{b \cdot d}$$

auf Rechenoperationen mit den natürlichen Zahlen a, b, c, d zurückgeführt werden!

2. Es ist zu beweisen, daß die Addition und die Multiplikation von Brüchen den sieben Grundgesetzen der Addition und Multiplikation natürlicher Zahlen genügen.
3. An konkreten Beispielen ist zu bestätigen, daß die in der Aufgabe 1 angegebenen Regeln der Bruchrechnung davon unabhängig sind, ob die Brüche $\frac{a}{b}$ und $\frac{c}{d}$ soweit als möglich gekürzt sind oder nicht.

3. Der Bereich der rationalen Zahlen

Wir haben im vorigen Abschnitt die Erweiterung der natürlichen Zahlen zu den Brüchen mit natürlichen Zahlen als Zähler und Nenner besprochen. In dem so entstandenen Bereiche gelten die gleichen Rechengesetze wie im Bereiche der natürlichen Zahlen, darüber hinaus ist die Division uneingeschränkt durchführbar. Jedoch ist nach wie vor die Subtraktion nur dann möglich, wenn der Minuend größer als der Subtrahend ist. Diese Tatsache veranlaßt uns, auch diesen Bereich zu erweitern. Wir werden den **Bereich der rationalen Zahlen** erhalten.

Diese Erweiterung kann nicht einfach dadurch vorgenommen werden, daß wir den mit den Bildern der bereits vorhandenen Brüche markierten Zahlenstrahl (Abb. 36) an dem mit 0 bezeichneten Anfangspunkt spiegeln und alle neu ent-

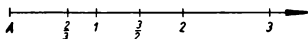


Abb. 36

stehenden Zahlen mit einem Minuszeichen versehen (Abb. 37). Die Punkte der so entstehenden Zahlengeraden sind nämlich nur Bilder der rationalen Zahlen und nicht diese selbst. Wir können aus diesem Bild nur entnehmen, wie die rationalen Zahlen angeordnet sind, nicht aber, wie mit ihnen zu rechnen ist. Man muß die Erweiterung vielmehr in der im vorigen Abschnitt besprochenen Art durchführen. Wir wollen dies in großen Zügen andeuten.

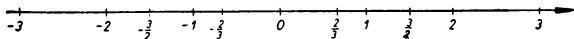


Abb. 37

Wir erklären die neuen (rationalen) Zahlen als die Gesamtheit aller Subtraktionsaufgaben mit alten Zahlen, also mit den bereits vorhandenen Brüchen. Eine Reihe von solchen Subtraktionsaufgaben fallen immer zu der gleichen neuen Zahl zusammen, zum Beispiel

$$3 - 5 = \frac{7}{2} - \frac{11}{2} = \frac{11}{3} - \frac{17}{3} = \frac{3}{7} - \frac{17}{7} = \dots$$

oder

$$1 - \frac{4}{3} = \frac{2}{3} - 1 = \frac{5}{6} - \frac{7}{6} = \frac{125}{3} - \frac{126}{3} = \dots$$

oder

$$7 - 2 = \frac{21}{4} - \frac{1}{4} = 6 - 1 = \frac{13}{2} - \frac{3}{2} = \dots$$

118 Vom Bereich der natürlichen Zahlen zum Bereich der reellen Zahlen

Man nennt dann die Zahlen, deren Minuend größer als der Subtrahend ist, positiv; diejenigen, deren Minuend kleiner als der Subtrahend ist, negativ, und führt kürzere Bezeichnungen ein (in unseren Beispielen -2 , $-\frac{1}{3}$, bzw. $+5$; vgl. Aufgabe 2).

Das Rechnen mit den neuen Zahlen muß dann auf das Rechnen mit den alten Zahlen zurückgeführt werden. So wird zum Beispiel die Aufgabe

$$(+5) + (-2) = (+3)$$

abgeleitet aus

$$[7 - 2] + [3 - 5] = [(7 + 3) - (2 + 5)] = (10 - 7),$$

wobei man die Minuenden 7 und 3 und die Subtrahenden 2 und 5 für sich addiert. Es läßt sich nachweisen, daß dies tatsächlich allgemein möglich ist und dabei alle für die Brüche gültigen Rechenregeln auch für die neuen Zahlen gültig bleiben (Permanenzprinzip). Die so entstehenden neuen Zahlen enthalten die alten als positive Zahlen und ermöglichen es, die Subtraktion uneingeschränkt auszuführen.

Damit haben wir in zwei Schritten einen Zahlbereich gewonnen, in dem ebenso gerechnet wird wie im Bereich der natürlichen Zahlen und in dem alle vier Grundrechenarten uneingeschränkt durchführbar sind, allerdings mit Ausnahme der Division durch Null. Man nennt diesen neuen Bereich den der rationalen Zahlen.

Wir wollen noch kurz begründen, daß die Aufgabe, durch Null zu dividieren, grundsätzlich undurchführbar ist. Der Quotient $\frac{a}{b}$ ist als die Lösung x der Gleichung $bx = a$ erklärt. Setzen wir aber in dieser Gleichung $b = 0$, so wird das Produkt bx unabhängig von der Wahl von x zu Null. Für $a \neq 0$ ist das ein Widerspruch, für $a = 0$ hätte die Gleichung alle Zahlen x als Lösung. Wir können also auch keine Erweiterung finden, die eine Lösung der Aufgabe $a : 0$ ermöglichte, ohne die für Zahlen gültigen Rechengesetze (Permanenzprinzip) zu verletzen.

Man kann noch auf einem zweiten Wege vom Bereich der natürlichen Zahlen zum Bereich der rationalen Zahlen gelangen. Anstatt zuerst die uneingeschränkte Ausführbarkeit der Division zu fordern und damit als Zwischenbereich den der positiven Brüche zu erhalten, kann man mit der Forderung der uneingeschränkten Durchführbarkeit der Subtraktion beginnen. Dann erhält man als Zwischenbereich den der **ganzen Zahlen**, das heißt den der positiven und negativen ganzen Zahlen einschließlich der Null. Von ihm gelangt man durch die Forderung der uneingeschränkten Durchführbarkeit der Division (außer der durch Null) ebenfalls zum Bereich der rationalen Zahlen. Während der von uns beschrittene Weg der historische ist, der wegen seiner besseren Anschaulichkeit auch im Unterricht gewählt wird, hat der zweite Weg gewisse Vorteile bei einer systematisierenden Übersicht. Es ist nämlich folgerichtiger, zuerst die Umkehrung der Addition als der Rechenoperation erster Stufe und dann die Umkehrung der Multiplikation als der Rechenoperation zweiter Stufe zu behandeln.

Aufgaben

1. Die Erweiterung des Bereiches der positiven Brüche zum Bereich der rationalen Zahlen ist schrittweise mit der Erweiterung des Bereiches der natürlichen Zahlen zum Bereich der positiven Brüche zu vergleichen.
2. Welche Subtraktionsaufgaben mit alten Zahlen (positive Brüche) erklären die neue Zahl Null?

4. Der Bereich der reellen Zahlen

Wir wollen nunmehr den Bereich der rationalen Zahlen erweitern. Dazu gehen wir nochmals zu den natürlichen Zahlen zurück, die, wie wir erkannt hatten, in bestimmter Reihenfolge angeordnet sind (Abb. 38). Die Sonderrolle der 1 wird aufgehoben, wenn wir den Bereich der natürlichen Zahlen durch die Forderung der uneingeschränkten Durchführbarkeit der Subtraktion zunächst zum Bereich der ganzen Zahlen erweitern (Abb. 39).

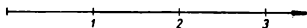


Abb. 38

Hier hat jede Zahl genau einen unmittelbaren Vorgänger und genau einen unmittelbaren Nachfolger. Daraus folgt, daß für zwei Zahlen genau eine der drei Beziehungen $a < b$, $a = b$, $a > b$ zutrifft.

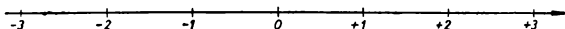


Abb. 39

Der Begriff der linearen Anordnung muß für die rationalen Zahlen etwas abgewandelt werden. Wir erinnern uns, daß jede rationale Zahl entweder als endlicher oder als unendlich-periodischer Dezimalbruch geschrieben werden kann (vgl. Lehrbuch der Mathematik 9. Schuljahr, S. 71). Teilen wir nun jeden Abschnitt der vorliegenden Zahlengeraden in 10 gleiche Teile, teilen diese Teilabschnitte ebenfalls in 10 gleiche Unterteile und denken dieses Verfahren soweit als irgend notwendig fortgesetzt (Abb. 40), so können wir jedem Dezimalbruch und damit jeder rationalen

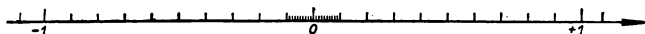


Abb. 40

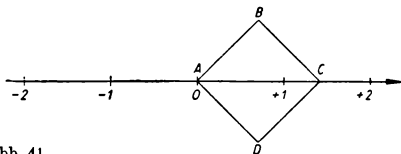
Zahl einen ganz bestimmten Punkt dieser so erweiterten Zahlengeraden zuordnen. In der damit entstehenden linearen Anordnung der rationalen Zahlen gibt es zwar keine unmittelbaren Vorgänger bzw. Nachfolger mehr, aber von zwei beliebigen rationalen Zahlen können wir an Hand ihrer Schreibweise als Dezimalbrüche feststellen, daß sie entweder gleich sind oder welche von beiden die größere, welche die kleinere ist. Es gilt also auch für rationale Zahlen genau eine der drei Beziehungen $a < b$, $a = b$, $a > b$ (vgl. Aufgabe 1).

Die Punkte, die die ganzen Zahlen darstellen, sind voneinander getrennt, sie liegen diskret. Zwischen zwei noch so benachbarten rationalen Zahlen a und b liegen dagegen immer noch unendlich viele weitere rationale Zahlen. Man sagt dazu auch, die den rationalen Zahlen entsprechenden Punkte liegen überall dicht. Auf den Beweis wollen wir verzichten.

Dennoch gibt es zwischen diesen überall dicht liegenden rationalen Punkten noch weitere Punkte, denen keine rationale Zahl entspricht. Als Beispiel wollen wir in der angegebenen Weise ein Quadrat $ABCD$ von der Seitenlänge 1 über der Zahlengeraden so konstruieren, daß die Diagonale AC mit der Zahlengeraden zusammenfällt (Abb. 41). Nach Konstruktion ist der Eckpunkt C ein Punkt der Zahlengeraden, und sein Abstand vom Punkte $A = 0$ ist nach dem Lehrsatz des Pytha-

goras $\sqrt{2}$. Würde nun dem Punkt C eine rationale Zahl entsprechen, so wäre $\sqrt{2}$ eine rationale Zahl. Dies ist jedoch nicht der Fall (vgl. Lehrbuch der Mathematik 9. Schuljahr). Wir wissen auch, daß dieser Punkt C keine Ausnahme ist.

Abb. 41



Die weitaus meisten Wurzeln und Logarithmen rationaler Zahlen sind nicht rational.

Andererseits liegt der Punkt C , gerade weil die rationalen Punkte überall dicht liegen, in unmittelbarer Nachbarschaft von solchen rationalen Punkten, das heißt, die Zahl $\sqrt{2}$ kann durch rationale Zahlen mit beliebiger Genauigkeit angenähert werden. Zum Beispiel gilt

$$\begin{aligned} 1 &< \sqrt{2} < 2 \\ 1,4 &< \sqrt{2} < 1,5 \\ 1,41 &< \sqrt{2} < 1,42 \\ 1,414 &< \sqrt{2} < 1,415 \\ &\vdots \end{aligned}$$

usw.

Dies lehrt uns zweierlei: Einmal können wir die Zahl $\sqrt{2}$ bei allen praktischen Rechnungen so genau, wie es nur erforderlich ist, durch rationale Zahlen annähern. Auf diese Weise haben wir bisher stets mit Wurzeln und Logarithmen gerechnet. Zum anderen aber erhalten wir einen Hinweis, wie wir den Bereich der rationalen Zahlen zu erweitern haben, um all die Aufgaben exakt lösen zu können, für die wir bisher nur Näherungslösungen ermitteln konnten. Die Tabelle zeigt uns nämlich, daß die Zahl $\sqrt{2}$ nichts anderes ist als der Grenzwert der Folge rationaler Zahlen

$$1; 1,4; 1,41; 1,414; \dots$$

Diese Folge konvergiert, sie hat aber keinen rationalen Grenzwert. Genauso verhält es sich bei allen anderen Annäherungen nicht rationaler Zahlen durch rationale. Wir werden also als neue Zahlen einfach die Gesamtheit aller konvergenten Folgen mit rationalen Zahlen einführen. Sie entspricht der Gesamtheit aller endlichen und unendlichen Dezimalbrüche und wird der **Bereich der reellen Zahlen** genannt.

Eine exakte Durchführung dieser Erweiterung müßte wieder so vorgenommen werden: Bekannt sind die rationalen Zahlen. Die reellen Zahlen werden als neue Zahlen mit Hilfe konvergenter Folgen rationaler Zahlen erklärt. Zwei Folgen, bei denen die Differenz der Glieder gegen Null konvergiert, definieren dabei die gleiche reelle Zahl. Das Rechnen mit diesen Folgen muß auf das Rechnen mit ihren Gliedern, also auf das bekannte Rechnen mit rationalen Zahlen zurückgeführt werden. Dabei müssen alle für rationale Zahlen geltenden Rechengesetze erhalten bleiben

(Permanenzprinzip). Jede rationale Zahl tritt selbst als eine solche Folge auf. Zum Beispiel wird die Zahl $\frac{1}{2}$ dargestellt durch die Folge

$$\frac{1}{2}; \quad \frac{1}{2}; \quad \frac{1}{2}; \quad \frac{1}{2};$$

oder durch die Folge

$$0; 0,4; 0,49; 0,499; 0,4999; \dots$$

Der auf diese Weise entstandene Bereich der reellen Zahlen ist damit eine Erweiterung des Bereiches der rationalen Zahlen. Die neuen, nicht rationalen reellen Zahlen werden **irrationale Zahlen** genannt.

Im Bereich der reellen Zahlen hat dann jede konvergierende Folge eine reelle Zahl als Grenzwert. Das bedeutet, daß sämtliche Wurzeln mit positivem Radikanden und sämtliche Logarithmen positiver Zahlen mit positiver Basis in diesem Bereich enthalten sind. Darüber hinaus enthält er auch diejenigen Zahlen, die überhaupt nur als Grenzwert einer Folge erklärt werden können, wie zum Beispiel die bei der Kreisberechnung auftretende Zahl π oder die Zahl e , die Basis der natürlichen Logarithmen.

Auch jeder irrationalen reellen Zahl entspricht ein bestimmter Punkt der Zahlengeraden. Auf Grund ihrer Konstruktion erschöpfen sie alle noch vorhandenen Punkte. Auf diese Weise entspricht jeder reellen Zahl genau ein Punkt der Zahlengeraden und umgekehrt jedem Punkt der Zahlengeraden genau eine rationale bzw. irrationale reelle Zahl. Damit ist auch für die reellen Zahlen wieder eine lineare Anordnung gegeben, wie wir sie schon für die rationalen Zahlen kennengelernt haben.

Aufgaben

1. Erklären Sie die Beziehungen $a < b$, $a = b$, $a > b$ zwischen Dezimalbrüchen unabhängig von ihrer Darstellung auf der Zahlengeraden nur mit Hilfe der in ihnen auftretenden Ziffern!
2. Geben Sie an, welche konvergierende Folge rationaler Zahlen und welcher unendliche Dezimalbruch die Irrationalzahl $\sqrt{1,2}$ darstellt! Es sind die ersten Glieder bzw. Ziffern anzugeben. Mit Hilfe einer geeigneten Reihendarstellung ist entsprechend mit $\ln 0,9$ zu verfahren.

II. Der Bereich der komplexen Zahlen

Ausgehend vom Bereich der natürlichen Zahlen sind wir durch mehrfache Erweiterung zum Bereich der reellen Zahlen gelangt. Auch in diesem Bereich sind noch nicht sämtliche Rechenoperationen uneingeschränkt durchführbar. Es ist zum Beispiel nicht möglich, die Quadratwurzel aus einer negativen Zahl auszu ziehen, da es keine reelle Zahl gibt, deren Quadrat negativ ist. Auch die Logarithmen negativer Zahlen existieren in diesem Bereich nicht.

Es liegt also nahe, auch über den Bereich der reellen Zahlen hinauszugehen und einen neuen Zahlenbereich zu schaffen, in dem alle sieben Rechenoperationen uneingeschränkt ausführbar sind. Da, wie wir gesehen haben, die reellen Zahlen die Punkte der Zahlengeraden bereits erschöpfen, können die Bilder der neuen Zahlen nicht mehr auf dieser Geraden liegen. Damit entfällt auch die Möglichkeit, diese Zahlen zusammen mit den reellen Zahlen linear anzuordnen. Wir sehen also, daß wir den ursprünglichen Begriff der (natürlichen) Zahl noch wesentlich mehr als bisher verallgemeinern müssen.

Aus dem Gesagten ist zu erklären, daß die historische Entwicklung verhältnismäßig lange bei den reellen Zahlen stehengeblieben ist. Bis zum Beginn der Neuzeit wurden Lösungen von Aufgaben zum Beispiel in der Form $x = \sqrt{-2}$ für unmöglich, unwirklich oder eingebildet gehalten. Dadurch ist auch die heute noch gebräuchliche Bezeichnung „imaginäre Zahlen“ für derartige Lösungen entstanden, obgleich wir dabei heute nicht mehr an die ursprüngliche Bedeutung des Wortes imaginär denken. Diese Zahlen sind zusammen mit den aus ihnen entwickelten komplexen Zahlen zu einem unentbehrlichen Hilfsmittel der gesamten Mathematik und ihrer Anwendungsgebiete geworden. Sie sind damit weder geheimnisvoller noch unwirklicher als die Zahlen der bisher betrachteten Bereiche.

5. Einführung der imaginären Zahlen

Wir wollen als Erweiterung des Bereiches der reellen Zahlen einen neuen Zahlenbereich finden, in dem die Rechengesetze der reellen Zahlen weitergelten und alle sieben Rechenoperationen (mit Ausnahme der Division durch Null) uneingeschränkt durchführbar sind. Dazu gehen wir schrittweise vor und betrachten zunächst die im Bereich der reellen Zahlen unlösbaren Quadratwurzeln, etwa

$$\sqrt{-1}, \quad \sqrt{-4}, \quad \sqrt{-\frac{5}{9}}, \quad \sqrt{-0,333 \dots}, \quad \sqrt{-\pi}.$$

Sie alle sind von der Form $\sqrt{-a}$ mit irgendeiner positiven reellen Zahl a . Diese im alten Bereich unlösbaren Rechenaufgaben führen wir als neue Zahlen ein; wir nennen sie **imaginäre Zahlen**. Ein solches Vorgehen entspricht genau unserem Vorgehen bei jeder der bisher besprochenen Erweiterung von Zahlenbereichen. Die imaginären Zahlen entstehen also auf genau die gleiche Art wie zum Beispiel die negativen oder irrationalen Zahlen. Wir müssen noch wie bei den bisher betrachteten Erweiterungen feststellen, wie mit diesen imaginären Zahlen umzugehen, das heißt zu rechnen ist.

Nach dem Permanenzprinzip sollen auch im neuen Zahlenbereich die Rechengesetze der reellen Zahlen gelten. Wir können daher wegen

$$\sqrt{a \cdot b} = \sqrt{a} \cdot \sqrt{b}$$

alle imaginären Zahlen $\sqrt{-a}$ mit $a > 0$ umformen¹⁾:

$$\sqrt{-a} = \sqrt{(+a) \cdot (-1)} = \sqrt{+a} \cdot \sqrt{-1}.$$

Da $\sqrt{+a}$ stets zwei reelle Lösungen hat, werden damit alle imaginären Zahlen zu reellzahligen Vielfachen von $\sqrt{-1}$, das heißt, alle im Bereich der reellen Zahlen unlösbaren Aufgaben $\sqrt{-a}$ werden auf eine, nämlich $\sqrt{-1}$, zurückgeführt.

Eine Lösung der Aufgabe $\sqrt{-1}$ (von denen der neue Bereich ja wenigstens eine enthalten muß) sei i . Mit

$$i^2 = -1$$

ist aber auch

$$(-i)^2 = (-1)^2 \cdot i^2 = -1.$$

Wir schreiben also

$$\sqrt{-1} = \pm i$$

und nennen i die imaginäre Einheit. Damit sind alle imaginären Zahlen reellzahlige Vielfache der imaginären Einheit i .

Aufgaben

1. Die folgenden imaginären Zahlen sind nach dem Beispiel

$$\sqrt{-4} = \sqrt{+4} \cdot \sqrt{-1} = \pm 2 \cdot i$$

als reellzahlige Vielfache der imaginären Einheit i zu schreiben:

$$\sqrt{-9}, \quad \sqrt{-7}, \quad \sqrt{-1}, \quad \sqrt{-\frac{5}{9}}, \quad \sqrt{-\frac{16}{25}}, \quad \sqrt{-\pi}, \quad \sqrt{-18\pi}.$$

2. Die gleiche Umformung wie in Aufgabe 1 ist für $\sqrt{-a^2}$ durchzuführen und das Ergebnis zu erläutern.

3. Lösen Sie die folgenden rein-quadratischen Gleichungen:

$$x^2 + 16 = 0, \quad x^2 + 3 = 0, \quad x^2 + 1 = 0, \quad x^2 - 16 = 0,$$

$$x^2 + 9 = 0, \quad x^2 + \frac{1}{7} = 0, \quad x^2 - 5 = 0, \quad x^2 = 0!$$

Was kann man über die Anzahl der Wurzeln jeder dieser Gleichungen aussagen?

4. Die folgenden reellen bzw. imaginären Zahlen sind als Quadratwurzeln zu schreiben:

$$\begin{array}{cccccc} +2, & -2, & +2 \cdot i, & -2 \cdot i, & -7, & +3 \cdot i, \\ -\pi \cdot i, & -i, & +a, & -a, & +a \cdot i, & -a \cdot i. \end{array}$$

(a ist eine beliebige reelle Zahl.)

5. Die Gleichung $x^4 - 16 = 0$ hat zwei reelle und zwei imaginäre Wurzeln. Geben Sie diese an und bilden Sie ähnliche Beispiele!

Anleitung: Es ist $\sqrt[4]{a} = \sqrt{\sqrt{a}}$.

¹⁾ Beachten Sie, daß $\sqrt{a}$ eine Rechenanweisung ist, eine abgekürzte Schreibweise für die Lösungen der Gleichung $x^2 = a$, von denen bekanntlich bei $a > 0$ im Bereich der reellen Zahlen stets genau zwei existieren! Desgleichen ist die Formel $\sqrt{a \cdot b} = \sqrt{a} \cdot \sqrt{b}$ als eine Relation zwischen diesen Rechenanweisungen aufzufassen. Beim Einsetzen konkreter Zahlen ist daher auf die Doppeldeutigkeit der Lösungen (Vorzeichen) zu achten.

6. Das Rechnen mit imaginären Zahlen

Wir wollen zunächst imaginäre Zahlen miteinander und mit reellen Zahlen multiplizieren. Da nach dem Permanenzprinzip die gleichen formalen Rechengesetze wie für reelle Zahlen gelten sollen, macht dies keinerlei Schwierigkeiten. Man muß nur berücksichtigen, daß wegen

$$\sqrt{-1} = \pm i$$

stets für $i^2 = (\sqrt{-1})^2 = -1$ zu setzen ist.

Es ist also

oder allgemein

$$i^1 = i$$

$$i^2 = -1$$

$$i^3 = i^2 \cdot i = (-1) \cdot i = -i$$

$$i^4 = i^2 \cdot i^2 = (-1) \cdot (-1) = +1$$

$$i^5 = i^4 \cdot i = +1 \cdot i = i \quad \text{usw.}$$

$$i^{4n} = +1$$

$$i^{4n+1} = +i$$

$$i^{4n+2} = -1$$

$$i^{4n+3} = -i$$

mit $n = 0; 1; 2; \dots$

Mit Hilfe dieser Gesetzmäßigkeit kann man die folgenden Produkte ausrechnen:

$$3 \cdot (5 \cdot i) = (3 \cdot 5) \cdot i = 15 \cdot i$$

$$(3 \cdot i) \cdot (5 \cdot i) = (3 \cdot 5) \cdot i^2 = 15 \cdot (-1) = -15$$

$$(3 \cdot i) \cdot (5 \cdot i) \cdot (2 \cdot i) = (3 \cdot 5 \cdot 2) \cdot i^3 = 30 \cdot (-i) = -30 \cdot i$$

$$a \cdot (b \cdot i) = (a \cdot b) \cdot i$$

$$(a \cdot i) \cdot (b \cdot i) = (a \cdot b) \cdot i^2 = -(a \cdot b).$$

Zur Vereinfachung der Schreibweise werden wir im folgenden den Multiplikationspunkt vor der imaginären Einheit i meist weglassen, also statt $3 \cdot i$ einfach $3i$, statt $a \cdot i$ einfach ai schreiben.

Als nächstes wollen wir Divisionsaufgaben mit reellen und imaginären Zahlen lösen. Auch da ergeben sich wegen der Permanenz der Rechengesetze aus

$$i \cdot i = -1 \quad \text{die Formeln} \quad i = \frac{-1}{i} \quad \text{und} \quad -i = \frac{1}{i}.$$

Wir können also die folgenden Brüche umformen:

$$\frac{15i}{3} = \frac{15}{3} i = 5i,$$

$$\frac{7}{3i} = \frac{7}{3} \cdot \frac{1}{i} = \frac{7}{3} (-i) = -\frac{7}{3} i,$$

$$\frac{8i}{4i} = \frac{8}{4} \cdot \frac{i}{i} = 2 \cdot 1 = 2.$$

Bei der Addition und Subtraktion wollen wir uns zunächst auf solche Aufgaben beschränken, die nur imaginäre Zahlen enthalten. Man kann, da auch im neuen Bereich das distributive Gesetz gelten soll, die imaginäre Einheit ausklammern. So gilt zum Beispiel:

$$\begin{aligned}5i + 7i &= 5 \cdot i + 7 \cdot i = (5 + 7) \cdot i = 12i, \\17i - 19i &= 17 \cdot i - 19 \cdot i = (17 - 19) \cdot i = -2i.\end{aligned}$$

Allgemein ist

$$ai + bi = (a + b)i.$$

Aufgaben

1. Die folgenden Aufgaben sind zu lösen:

$$\begin{aligned}\text{a) } 3i \cdot 16; & \quad 10i \cdot 7 \cdot 5i; & \quad 16 \cdot \frac{5}{2}i; & \quad \frac{1}{3} \cdot \pi \cdot i; & \quad 5 \cdot \sqrt{-25}; \\3i \cdot \sqrt{-8}; & \quad \sqrt{-12} \cdot \sqrt{-\frac{1}{4}}; & \quad \frac{1}{30} \cdot \frac{15}{3}i \cdot \sqrt{-8}; & \quad (\sqrt{-5})i \cdot (\sqrt{-9})i.\end{aligned}$$

b) Welche Grundgesetze der Multiplikation werden zur Lösung der Aufgaben verwendet?

2. Begründen Sie, weshalb das Produkt zweier reeller Zahlen reell, das Produkt einer reellen mit einer imaginären Zahl imaginär, das Produkt zweier imaginärer Zahlen reell ist!

3. Stellen Sie eine Tabelle auf, die die Potenzen von i mit negativen Exponenten von i^{-1} bis i^{-8} enthält!

4. Die folgenden Brüche sind zu vereinfachen:

$$\frac{3}{7i}; \quad \frac{\sqrt{-4}}{\sqrt{3}}; \quad \frac{17}{\sqrt{-7}}; \quad \frac{10i}{\sqrt{-2}} \cdot 5i; \quad \frac{1}{\sqrt{-\pi}} \cdot \frac{3\pi}{2i}; \quad \frac{a}{bi}; \quad \frac{ai}{b}; \quad \frac{ai}{bi}.$$

5. Es ist zu berechnen:

$$\begin{aligned}\frac{5}{3}i - \frac{3}{2}i + \frac{5}{6}i; & \quad 15i - \sqrt{-3}; & \quad \frac{\sqrt{-2} + 7i}{15i}; & \quad \frac{3i}{5i} + \frac{\sqrt{-4} + \sqrt{-9}}{25i}; \\ \sqrt{-\frac{1}{4}} + 4i - \sqrt{-\frac{3}{8}}; & \quad 5 \cdot \sqrt{-\frac{4}{3}} + \frac{10}{3i}.\end{aligned}$$

6. Berechnen Sie

$$\frac{4}{7}i - \frac{3}{4}i + \sqrt{-2}; \quad \sqrt{-\frac{1}{4}} - 3i - \sqrt{-\frac{1}{9}}; \quad 4 \cdot \sqrt{-\frac{4}{9}} - \frac{5}{6i} + \frac{6}{9}i!$$

7. Wodurch unterscheidet sich die Aufgabe

$$17i + \frac{3i}{\sqrt{-5}}$$

von den unter Nummer 5 angegebenen Aufgaben?

7. Einführung der komplexen Zahlen

Wir müssen nun noch untersuchen, zu welchem Ergebnis wir bei der Addition bzw. Subtraktion von reellen und imaginären Zahlen gelangen. Auch diese Rechenoperationen müssen sich in dem angestrebten neuen Zahlenbereich ausführen lassen. Versuchen wir aber, eine reelle Zahl und eine imaginäre Zahl zu addieren, etwa

$$3 + 5i$$

oder allgemein

$$a + bi \quad (a \neq 0; \quad b \neq 0),$$

so kann das Ergebnis weder eine reelle noch eine imaginäre Zahl sein.

Wäre nämlich $a + bi = c$ mit irgendeiner geeigneten reellen Zahl c , so folgte nach dem Permanenzprinzip, daß

$$bi = c - a,$$

also reell sein müßte, was doch nicht der Fall ist.

Analog folgte aus $a + bi = di$, daß a imaginär sein müßte.

Damit ergibt sich eine weitere Notwendigkeit: Es genügt nicht, den Bereich der reellen Zahlen nur mit den zum Quadratwurzelausziehen notwendigen imaginären Zahlen zu erweitern. Wir müssen vielmehr auch alle Summen

$$a + bi$$

einer beliebigen reellen Zahl a und einer beliebigen imaginären Zahl bi als von einander verschiedene Zahlen des neuen Bereiches ansehen. Ihrer zusammengesetzten Form wegen werden sie **komplexe Zahlen** genannt, wobei man die reelle Zahl a als den **Realteil**, die reelle Zahl b als den **Imaginärteil** der komplexen Zahl $a + bi$ bezeichnet. Die Gesamtheit aller komplexen Zahlen, das heißt also aller Ausdrücke $a + bi$ mit beliebigem reellem a und beliebigem reellem b , bildet bereits den von uns gesuchten neuen Zahlenbereich. Er enthält insbesondere auch alle reellen und alle imaginären Zahlen. Für reelle Zahlen ist nämlich $b = 0$ und für imaginäre $a = 0$. Wir werden sehen, daß in diesem Bereich alle sieben Grundrechenoperationen uneingeschränkt ausführbar sind.

Aufgaben

1. Es ist Realteil und Imaginärteil der folgenden komplexen Zahlen zu nennen:

$$\begin{array}{lllll} 3 + 5i; & 4 + i; & -5 + 2i; & \frac{1}{3} + (-4)i; & \frac{1}{3} - 4i; \\ -12 - 8i; & 1 + i; & x - yi; & 0 + 3i; & \frac{1}{2} + 0i. \end{array}$$

2. Das Ergebnis der Aufgabe 7 im vorigen Abschnitt ist eine komplexe Zahl. Geben Sie ihren Realteil und ihren Imaginärteil an!

8. Die folgenden Zahlen sind als komplexe Zahlen zu schreiben:

$$17; \quad 17i; \quad 0,33i; \quad \sqrt{-5}; \quad \frac{538}{5}; \quad \frac{3i}{\sqrt{-4}}; \quad -1; \quad 0.$$

8. Die Rechenoperationen erster und zweiter Stufe im Bereich der komplexen Zahlen

Nach dem Permanenzprinzip sollen alle Rechengesetze, die wir für den Bereich der reellen Zahlen kennengelernt haben, auch für komplexe Zahlen gelten. Dazu müssen wir die Summe und das Produkt zweier komplexer Zahlen geeignet festsetzen, das heißt auf Rechnungen mit reellen Zahlen zurückführen. Wir werden sehen, daß man dabei unter Berücksichtigung von $i^2 = -1$ so verfahren kann, als ob alle vorkommenden Ausdrücke reell wären. Mit den auf diese Weise gewonnenen Erklärungen für Summe und Produkt gelten dann tatsächlich im Bereich der komplexen Zahlen die gleichen Rechengesetze wie im Bereich der reellen Zahlen.

Um eine geeignete Erklärung für die Summe zweier komplexer Zahlen

$$(a + bi) + (c + di)$$

zu finden, ersetzen wir die imaginäre Einheit i durch eine reelle Zahl x und erhalten

$$(a + bx) + (c + dx).$$

Dann gilt

$$(a + bx) + (c + dx) = a + c + bx + dx = (a + c) + (b + d)x.$$

Entsprechend setzen wir bei der ursprünglichen Aufgabe

$$(a + bi) + (c + di) = a + c + bi + di = (a + c) + (b + d)i$$

und erklären:

Komplexe Zahlen werden addiert, indem die Realteile und die Imaginärteile für sich addiert werden.

Damit ist die Summe zweier komplexer Zahlen stets wieder eine eindeutig bestimmte komplexe Zahl. Diese Formulierung ist auch dann richtig, wenn in besonderen Fällen das Ergebnis reell oder imaginär ausfällt. Wie wir im vorigen Abschnitt gesehen haben, sind ja die reellen Zahlen die komplexen Zahlen $a + bi$ mit $b = 0$, die imaginären Zahlen die komplexen Zahlen $a + bi$ mit $a = 0$.

Die Addition komplexer Zahlen ist kommutativ und assoziativ (vgl. Aufg. 3). Sie ist aber auch umkehrbar, denn zu zwei komplexen Zahlen

$$(a + bi) \quad \text{und} \quad (c + di)$$

gibt es stets eine Differenz

$$(a + bi) - (c + di) = (a - c) + (b - d)i,$$

die als komplexe Zahl der Gleichung

$$(c + di) + ((a - c) + (b - d)i) = (a + bi)$$

genügt.

Komplexe Zahlen werden also subtrahiert, indem man ebenfalls die Realteile und die Imaginärteile für sich subtrahiert.

Bevor wir zur Multiplikation übergehen, wollen wir noch eine Begriffsbildung kennenlernen, die beim Rechnen mit komplexen Zahlen häufig verwendet wird. Zu jeder komplexen Zahl $a + bi$ gibt es eine komplexe Zahl $a - bi$. Beide unterscheiden sich also nur durch das Vorzeichen der Imaginärteile und werden **konjugiert komplex** genannt. Es ist auch üblich, die eine Zahl den konjugiert komplexen Wert der anderen zu nennen.

Auf Grund der bisher besprochenen Gesetze für komplexe Zahlen gilt

$$(a + bi) + (a - bi) = 2a$$

und

$$(a + bi) - (a - bi) = 2bi.$$

Die Summe zweier konjugiert komplexer Zahlen ist also reell, die Differenz zweier konjugiert komplexer Zahlen imaginär.

Um eine geeignete Erklärung für das Produkt zweier komplexer Zahlen

$$(a + bi) \cdot (c + di)$$

zu finden, verfahren wir wie bei der Summe. Ersetzen wir wieder die imaginäre Einheit i durch die reelle Zahl x , so erhalten wir

$$\begin{aligned}(a + bx) \cdot (c + dx) &= ac + adx + bxc + bxdx \\ &= ac + (ad + bc)x + bdx^2.\end{aligned}$$

Entsprechend setzen wir bei der ursprünglichen Aufgabe

$$(a + bi) \cdot (c + di) = ac + (ad + bc)i + bdi^2.$$

Unter Berücksichtigung von $i^2 = -1$ erhalten wir

$$(a + bi) \cdot (c + di) = (ac - bd) + (ad + bc)i$$

und erklären:

Komplexe Zahlen werden nach der Formel

$$(a + bi) \cdot (c + di) = (ac - bd) + (ad + bc)i$$

multipliziert.

Damit ist das Produkt zweier komplexer Zahlen stets wieder eine eindeutig bestimmte komplexe Zahl. Insbesondere ist das Produkt zweier konjugiert komplexer Zahlen

$$(a + bi) \cdot (a - bi) = a^2 + b^2 + (ab - ba)i = a^2 + b^2$$

immer reell. Auf die Bedeutung von $a^2 + b^2$ werden wir später zu sprechen kommen.

Auch die Multiplikation komplexer Zahlen ist kommutativ und assoziativ. Dies läßt sich ohne Schwierigkeiten nachprüfen, doch wollen wir uns damit begnügen, beide Gesetze in Formeln aufzuschreiben (vgl. Aufgabe 9). Weiterhin ist die Multiplikation ebenfalls umkehrbar. Der Quotient zweier komplexer Zahlen

$$\frac{a + bi}{c + di}$$

ist nämlich wieder eine komplexe Zahl. Wir erhalten sie dadurch, daß wir mit dem konjugiert komplexen Wert des Nenners erweitern:

$$\frac{a + bi}{c + di} = \frac{(a + bi)(c - di)}{(c + di)(c - di)} = \frac{(ac + bd) + (bc - ad)i}{c^2 + d^2}.$$

Damit wird der Nenner $c^2 + d^2$ reell. Dividieren wir Realteil und Imaginärteil des Zählers für sich durch den jetzt reellen Nenner, so entsteht die komplexe Zahl

$$\frac{ac + bd}{c^2 + d^2} + \frac{bc - ad}{c^2 + d^2}i,$$

deren Produkt mit $c + di$ genau $a + bi$ ergibt (vgl. Aufgabe 11).

Dieser Übergang ist immer durchführbar, es sei denn, daß $c^2 + d^2 = 0$, das heißt also $c = d = 0$ ist (vgl. Aufgabe 12).

Im Bereich der komplexen Zahlen ist also jede Division, mit Ausnahme der durch 0, ausführbar.

Wir haben damit gelernt, sämtliche vier Rechenoperationen erster und zweiter Stufe mit komplexen Zahlen vorzunehmen. Sie genügen den bereits aus anderen Zahlenbereichen bekannten Rechengesetzen, denn außer den schon bei Addition und Multiplikation erwähnten Grundgesetzen gilt auch das beide Operationen verbindende distributive Gesetz (Aufgabe 13).

Wenn wir den Bereich der komplexen Zahlen mit dem Bereich der reellen Zahlen und dem Bereich der rationalen Zahlen vergleichen, erkennen wir: Solange wir nur die Rechenoperationen erster und zweiter Stufe betrachten, sind alle drei Zahlenbereiche bezüglich der dabei gültigen Rechengesetze gleich geartet. In jedem der drei Bereiche ist die Summe zweier Zahlen wieder eine bestimmte Zahl des Bereiches. Die Addition ist kommutativ und assoziativ und innerhalb jedes der drei Bereiche umkehrbar, das heißt, jede Subtraktionsaufgabe ist lösbar. Dies gilt sowohl für die Summe bzw. Differenz von rationalen Zahlen wie für die von reellen Zahlen als auch für die von komplexen Zahlen, wobei natürlich der Bereich der reellen Zahlen auch alle rationalen, der Bereich der komplexen Zahlen auch alle reellen Zahlen enthält.

Gleiches gilt für die Multiplikation in jedem der drei Bereiche. Sie ist stets eindeutig ausführbar, kommutativ, assoziativ und bis auf die Division durch Null umkehrbar. Schließlich gelten in jedem der drei Bereiche das distributive Gesetz und damit alle aus den Grundgesetzen folgenden Rechenregeln.

Anders wird es freilich, wenn wir die Rechenoperationen dritter Stufe betrachten. Zwar ist das Potenzieren mit einer natürlichen Zahl als Exponenten, das ja nichts anderes ist als eine Zusammenfassung von Multiplikationen, in allen drei Bereichen immer ausführbar und genügt den gleichen Gesetzen. Die beiden Umkehrope-

rationen sind dagegen im Bereich der rationalen Zahlen nur in speziellen Fällen, im Bereich der reellen Zahlen für alle positiven Zahlen, im Bereich der komplexen Zahlen stets ausführbar.

Zunächst wollen wir uns mit den komplexen Zahlen etwas vertrauter machen. Die reellen Zahlen konnten wir als Punkte einer Zahlengeraden veranschaulichen. Eine solche lineare Anordnung ist, wie wir bereits wissen, für komplexe Zahlen nicht mehr möglich, da die reellen Zahlen schon alle Punkte der Zahlengeraden erschöpfen. Die komplexen Zahlen können aber durch eine zweidimensionale Anordnung als Punkte einer Ebene veranschaulicht werden. Diese Veranschaulichung wurde von Carl Friedrich Gauß eingeführt. Sie heißt deshalb **Gaußsche Zahlenebene**.

Aufgaben

1. Berechnen Sie

- a) $(5 + 3i) + (7 + 2i)$; $(17 - 4i) + (3 + 2i)$; $(-8 + 4i) + (4 - 3i)$;
 b) $(0,75 + 2i) + (-0,5 - 0,5i)$; $3i + (4 - i)$; $(17 + 5i) + (4 - 5i)$;
 c) $(5 + 2i) + (-3 - 2i) + (4 - 7i)$; $(-4 + 3i) + 5i$; $5 + (17 - 2i)$;
 d) $(8 - 2i) + (2 + 2i)$; $(17 - i) + i$; $(8 + 3i) + (-8 + i)$;
 e) $(125 + 3i) + (-10 + i) + (2 - 4i)$; $(-4 - i) + (4 + 17i)$; $(-8 + 3i) + 8$;
 f) $(2 + 3i) + 17 + (-19 + 2i)$; $(a + bi) + (a - bi)$; $(a + bi) + (-a + bi)$!

2. Erklären Sie, wieso die Addition komplexer Zahlen auf Rechenoperationen mit reellen Zahlen zurückgeführt wird!

3. Geben Sie das kommutative und das assoziative Gesetz der Addition für komplexe Zahlen als Formeln an! Begründen Sie die Gültigkeit dieser Gesetze!

4. a) $(12 + 3i) - (4 + 2i)$ $(8 + 2i) - (3 + 6i)$ $(17 + 2i) - 9i$
 b) $(12 + 7i) - 8$ $(6 + 3i) - (9 - 2i)$ $(3 + 2i) - (-4 + i)$
 c) $(a + bi) - (a - bi)$ $(a + bi) - (-a + bi)$ $(a + bi) - (-a - bi)$

5. a) Zu den folgenden komplexen Zahlen sind die konjugiert komplexen Zahlen zu bilden.

$4 + 2i$	$1 - 2i$	$5 + i$	$0,5 - 0,75i$	$-4 + 2i$
$4 - 2i$	$a - bi$	$3a + 2bi$	$17i$	

- b) Aus den jeweils zueinander konjugiert komplexen Zahlen sind die Summe und die Differenz zu bilden.

6. Der folgende Satz ist zu begründen: Eine komplexe Zahl ist reell, wenn sie mit ihrer konjugiert komplexen Zahl zusammenfällt, und umgekehrt.

7. Die entsprechenden Betrachtungen wie bei Aufgabe 2 sind für die Multiplikation durchzuführen.

8. a) $(4 + 2i) \cdot (5 + 3i)$ $(4 - i) \cdot (3 + 3i)$ $(0,5i - 7) \cdot (8 - 2i)$
 b) $(3 + 2i) \cdot (3 - 2i)$ $(-5 + i) \cdot (-2 - 0,75i)$ $(2 - i) \cdot (2 + i)$
 c) $(-8 + 2,5i) \cdot (-8 - 2,5i)$ $(-1 - i) \cdot (-1 + i)$ $(-3 + 4i) \cdot (6 + 2i)$
 d) $(5 + 2i) \cdot (4 - 3i) \cdot (-7 - i)$ $(-5 - i) \cdot (1 + 3i) \cdot 2i$ $(2 - i) \cdot (9 + i)$

9. Schreiben Sie das kommutative und das assoziative Gesetz der Multiplikation für komplexe Zahlen als Formeln!

10. Berechnen Sie

- a) $\frac{10 + 10i}{2 + 4i}$; b) $\frac{1 - 2i}{2 + i}$; c) $\frac{-33 - 21i}{9 - 2i}$; d) $\frac{5 - 3i}{5 + 3i}$; e) $\frac{17 - i}{\sqrt{-9}}$;
 f) $\frac{\sqrt{-2}}{3 - 2i}$; g) $\frac{a + bi}{a - bi}$; h) $\frac{14xy + (4y^2 - 6x^2)i}{3x + yi}$

11. Es ist die im Text Seite 129, Zeile 12 aufgestellte Behauptung durch Ausrechnen nachzuweisen.

12. Bei dem im Text verwendeten Schluß, daß mit $c^2 + d^2 = 0$ auch $c = d = 0$ gilt, ist Voraussetzung, daß c und d als Realteil und Imaginärteil einer komplexen Zahl reell sind. An Beispielen ist zu zeigen, daß die Quadratsumme zweier komplexer Zahlen Null sein kann, ohne daß die beiden Zahlen Null sind.

13. Berechnen Sie

$$(3 + 2i) \cdot [(5 - i) + (1 + 3i)]$$

entsprechend dem distributiven Gesetz auf zwei verschiedene Arten und vergleichen Sie beide Ergebnisse. Es ist auch das distributive Gesetz für komplexe Zahlen als allgemeine Formel zu schreiben.

Entsprechend ist zu berechnen:

- a) $(7 - 4i) \cdot \left[\left(-\frac{1}{2} - \frac{1}{7}i \right) + \left(5 - \frac{3}{7}i \right) \right]$ b) $(5 + 2i) \cdot [(-3 + i) - (2 - 2i)]$
 c) $(-4 + 3i) \cdot [2i + (4 - 7i)]$ d) $9i \cdot \left[(5 + 2i) - \frac{1}{3}i \right]$
 e) $(\sqrt{4} + \sqrt{-4}) \cdot (8 - 5i) + \sqrt{-9}$ f) $(\sqrt{9} - \sqrt{-9}) \cdot (3 + 4i) - \sqrt{-4}$

9. Die Gaußsche Zahlenebene

Wie wir wissen, entspricht jedem Punkt der Zahlengeraden eine reelle Zahl und umgekehrt. Dabei ist diese Zuordnung schon durch die Wahl der Einheitsstrecke von 0 bis +1 festgelegt. Um dies einzusehen, brauchen wir nur jeder reellen Zahl den Vektor vom Ursprung 0 bis zu dem mit ihr bezeichneten Punkt zuzuordnen. Jeder reellen Zahl $a = a \cdot 1$ entspricht dann das a -fache des Einheitsvektors.

Entsprechend können wir alle imaginären Zahlen auf einer Geraden veranschaulichen, indem wir jede imaginäre Zahl $b \cdot i$ als das b -fache der imaginären Einheit i

auffassen. Die so aus der „imaginären Einheitsstrecke“ von 0 bis i entstehende Gerade wählen wir senkrecht zur Zahlengeraden mit dem Schnittpunkt $0 = 0 \cdot 1 = 0 \cdot i$ (vgl. Abb. 42). Die beiden Geraden werden reelle bzw. imaginäre

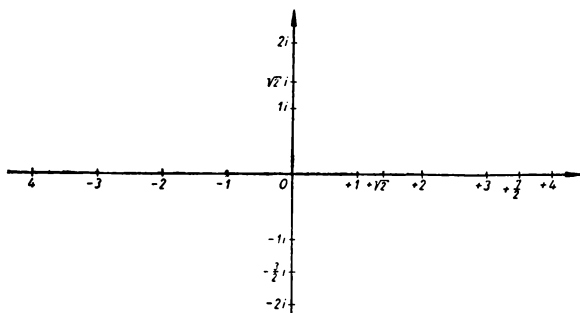


Abb. 42

Achse genannt. Die eine veranschaulicht alle reellen, die andere alle imaginären Zahlen. Die Zahl Null ist die einzige, die als reelle und zugleich als imaginäre Zahl aufgefaßt werden kann.

Aus der analytischen Geometrie ist bekannt, daß man jeden Punkt einer Ebene eindeutig durch ein Paar reeller Zahlen (Abszisse und Ordinate) kennzeichnen kann und daß umgekehrt jedem solchen Zahlenpaar eindeutig ein Punkt dieser Ebene entspricht. Andererseits ist jede komplexe Zahl

$$a + bi$$

durch die beiden reellen Zahlen a und b eindeutig festgelegt, und umgekehrt gibt es zu jedem Paar reeller Zahlen a und b genau eine komplexe Zahl $a + bi$. Beides zusammen bedeutet, daß durch die Konstruktion, wie sie die Abbildung 43 darstellt, jeder komplexen Zahl $a + bi$ genau ein Punkt der durch die reelle und imaginäre Achse aufgespannten Gaußschen Zahlenebene zugeordnet wird. Auf diese Weise entspricht also jedem Punkt der Zahlenebene eindeutig eine komplexe Zahl und jeder komplexen Zahl eindeutig ein Punkt der Zahlenebene. Die Punkte der reellen und der imaginären Zahlen fallen dabei auf die beiden Achsen. Die Abbildung 44 gibt dazu einige Beispiele.

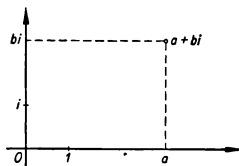


Abb. 43

Die Veranschaulichung der komplexen Zahlen als Punkte der Gaußschen Zahlenebene erweist sich auch für das praktische Rechnen mit ihnen oft als vorteilhaft. So können wir zum Beispiel die Addition komplexer Zahlen zeichnerisch aus-

führen, indem wir die Strecken wie Vektoren aneinanderfügen, die vom Nullpunkt zu den die komplexen Zahlen darstellenden Punkten führen (vgl. Abb. 45). Bei der Subtraktion kann man entsprechend verfahren (vgl. Abb. 46). Das darf uns jedoch nicht dazu veranlassen, die komplexen Zahlen

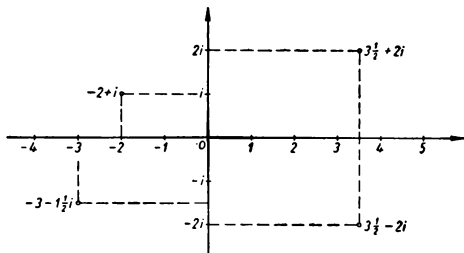


Abb. 44

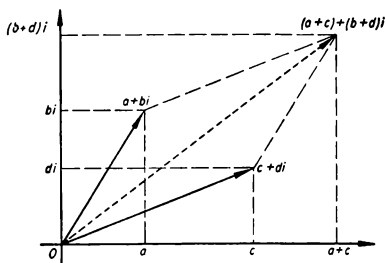


Abb. 45

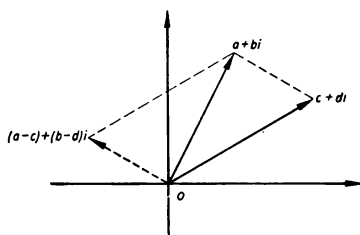


Abb. 46

selbst mit den Vektoren einer Ebene zu verwechseln. Komplexe Zahlen sind arithmetische oder algebraische Rechengrößen. Sie werden durch Punkte veranschaulicht, und die Verbindungsstrecken dieser Punkte zum Nullpunkt verhalten sich bezüglich Addition und Subtraktion wie Vektoren. Schon für die Multiplikation komplexer Zahlen gelten andere Gesetze als für die Multiplikation von Vektoren, so daß kein Vergleich möglich ist.

Aufgaben

1. Es ist zu untersuchen, welchen Punkten der Gaußschen Zahlenebene die folgenden komplexen Zahlen entsprechen.

$$3 + 4i$$

$$0,5 - 2i$$

$$-0,5 - 3i$$

$$-7 + 5i$$

$$1 + i$$

$$\sqrt[3]{5 \cdot i}$$

$$\sqrt{-8}$$

$$\pi$$

$$\pi i$$

$$0$$

2. Nach dem Vorbilde der analytischen Geometrie teilt man auch die Gaußsche Zahlenebene in vier Quadranten ein. Welche Vorzeichen haben Real- und Imaginärteil in diesen vier Quadranten?
3. Wo liegen alle komplexen Zahlen mit gleichem Realteil, wo die mit gleichem Imaginärteil?

4. Wie liegen konjugiert komplexe Zahlen zueinander?
5. Welchen Punkten der Gaußschen Zahlenebene entsprechen die komplexen Zahlen $a + bi$ mit $a = b$?
6. Wo liegen diejenigen komplexen Zahlen $a + bi$, für die die Gleichung $a^2 + b^2 = 1$ erfüllt ist?
7. Die folgenden Aufgaben sind zeichnerisch zu lösen und die Ergebnisse rechnerisch nachzuprüfen:
- | | | | |
|-----------------------|-----------------------|-----------------|-------------------------|
| $(3 + 2i) + (4 + i);$ | $(2 - i) - (3 + 3i);$ | $5i + (4 - i);$ | $(-3 + 2i) + (3 + 2i);$ |
| $(2 + i) + (2 - i);$ | $(2 + i) - (2 - i);$ | $(4 + 3i) - 6;$ | $6i - 5.$ |

III. Rechenoperationen im Bereich der komplexen Zahlen

10. Die trigonometrische Form der komplexen Zahlen

Wir haben die komplexen Zahlen in der sogenannten „arithmetischen Form“ $a + bi$ als alle möglichen Summen reeller und imaginärer Zahlen kennengelernt. Ihre Veranschaulichung in der Gaußschen Zahlenebene führt uns zu einer anderen Schreibweise, die die „trigonometrische Form“ der komplexen Zahlen genannt wird und die für die Durchführung der Rechenoperationen zweiter und dritter Stufe vorteilhafter ist.

Jeder komplexen Zahl $a + bi$ entspricht genau ein Punkt der Gaußschen Zahlenebene. Analytisch betrachtet sind dabei a und b die Cartesischen Koordinaten dieses Punktes. Gehen wir von ihnen zu Polarkoordinaten (vgl. Lehrbuch der Mathematik, 11. Schuljahr, Seite 191ff.) über, so erhalten wir zwei andere Bestimmungsstücke, r und φ , die diesen Punkt ebenfalls eindeutig kennzeichnen (vgl. Abb. 47).

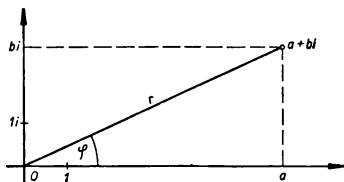


Abb. 47

Durch die Formeln¹⁾

$$r = +\sqrt{a^2 + b^2},$$

$$\varphi = \operatorname{arctg} \frac{b}{a}$$

und

$$a = r \cdot \cos \varphi,$$

$$b = r \cdot \sin \varphi$$

¹⁾ Das tiefgestellte Pluszeichen vor der Wurzel bedeutet, daß die Wurzel positiv auszuziehen ist.

lassen sich r und φ eindeutig aus a und b bestimmen und umgekehrt. Wir können daher die komplexe Zahl $a + bi$ durch r und φ ausdrücken:

$$a + bi = r \cdot \cos \varphi + i \cdot r \cdot \sin \varphi = r (\cos \varphi + i \cdot \sin \varphi).$$

Man nennt r den **Betrag**, φ das **Argument** und $r (\cos \varphi + i \cdot \sin \varphi)$ die **trigonometrische Form der komplexen Zahl**. Wegen der Periodizität der trigonometrischen Funktionen ist das Argument φ nur bis auf ganzzahlige Vielfache von 2π bestimmt. Der Betrag r ist dabei die positive Wurzel aus dem Produkt der komplexen Zahl mit ihrem konjugiert komplexen Wert:

$$r = +\sqrt{(a + bi) \cdot (a - bi)} = +\sqrt{a^2 + b^2},$$

wofür man auch

$$r = |a + bi|$$

schreibt. Diese Definition des Betrages einer komplexen Zahl steht nicht im Widerspruch zur Definition des Betrages einer reellen Zahl. Für $b = 0$ ist nämlich die komplexe Zahl $a + bi$ gleich der reellen Zahl a und ihr Betrag $r = |a + bi| = |a|$ der Betrag dieser reellen Zahl.

Aufgaben und Übungen

1. Welche komplexen Zahlen haben den gleichen Betrag, welche das gleiche Argument?
2. Welche Besonderheit gilt für die Null als komplexe Zahl in trigonometrischer Form?
3. Geben Sie die folgenden komplexen Zahlen jeweils in ihrer anderen Form an:

$$\text{a) } 3 + 4i; \quad \text{b) } \frac{1}{2} - \frac{1}{2}i \cdot \sqrt{3}; \quad \text{c) } +2 - 2i;$$

$$\text{d) } -\sqrt{3} - 3i; \quad \text{e) } 5i; \quad \text{f) } -1;$$

$$\text{g) } 2 \left(\cos \frac{\pi}{6} + i \cdot \sin \frac{\pi}{6} \right); \quad \text{h) } \sqrt{8} \left(\cos \frac{5\pi}{4} + i \cdot \sin \frac{5\pi}{4} \right);$$

$$\text{i) } 7 (\cos \pi + i \cdot \sin \pi); \quad \text{k) } \frac{2}{3} \left(\cos \frac{\pi}{3} + i \cdot \sin \frac{\pi}{3} \right);$$

$$\text{l) } \frac{2}{3} \left(\cos \frac{7\pi}{3} + i \cdot \sin \frac{7\pi}{3} \right); \quad \text{m) } 10 (\cos 53,1^\circ + i \cdot \sin 53,1^\circ);$$

$$\text{n) } 9 (\cos 270^\circ + i \cdot \sin 270^\circ)!$$

11. Multiplikation und Division komplexer Zahlen in trigonometrischer Form

Sind zwei komplexe Zahlen in arithmetischer Form gegeben, so lassen sie sich mühelos addieren und subtrahieren, während die Multiplikation und erst recht die Division einige Rechnungen erfordern. Diese Rechenoperationen lassen sich mit

komplexen Zahlen in trigonometrischer Form leichter durchführen. Es gilt nämlich, wie wir weiter unten beweisen,

$$\begin{aligned} \{r_1 (\cos \varphi_1 + i \cdot \sin \varphi_1)\} \cdot \{r_2 (\cos \varphi_2 + i \cdot \sin \varphi_2)\} \\ = r_1 \cdot r_2 (\cos (\varphi_1 + \varphi_2) + i \cdot \sin (\varphi_1 + \varphi_2)) \end{aligned}$$

und

$$\begin{aligned} \{r_1 (\cos \varphi_1 + i \cdot \sin \varphi_1)\} : \{r_2 (\cos \varphi_2 + i \cdot \sin \varphi_2)\} \\ = \frac{r_1}{r_2} (\cos (\varphi_1 - \varphi_2) + i \cdot \sin (\varphi_1 - \varphi_2)). \end{aligned}$$

Zwei komplexe Zahlen in trigonometrischer Form multipliziert man also, indem man ihre Beträge multipliziert und ihre Argumente addiert. Man dividiert sie, indem man ihre Beträge dividiert und ihre Argumente subtrahiert. Man muß nur beachten, daß sich bei der Addition bzw. Subtraktion der Argumente Werte ergeben können, die außerhalb des Winkelbereiches von 0 bis 2π liegen. Da negative Winkel oder Winkel größer als 2π auf Winkel im Bereich von 0 bis 2π zurückgeführt werden können, wird in einem solchen Falle das neue Argument um $+2\pi$ oder um -2π abgeändert.

Wir wollen nun die behaupteten Formeln beweisen. Zunächst ist

$$\begin{aligned} \{r_1 (\cos \varphi_1 + i \cdot \sin \varphi_1)\} \cdot \{r_2 (\cos \varphi_2 + i \cdot \sin \varphi_2)\} \\ = r_1 \cdot r_2 ([\cos \varphi_1 \cos \varphi_2 - \sin \varphi_1 \sin \varphi_2] + i [\sin \varphi_1 \cos \varphi_2 + \cos \varphi_1 \sin \varphi_2]). \end{aligned}$$

Nach den Additionstheoremen für trigonometrische Funktionen ist aber

$$\cos \varphi_1 \cos \varphi_2 - \sin \varphi_1 \sin \varphi_2 = \cos (\varphi_1 + \varphi_2)$$

und

$$\sin \varphi_1 \cos \varphi_2 + \cos \varphi_1 \sin \varphi_2 = \sin (\varphi_1 + \varphi_2),$$

woraus bereits die behauptete Formel für das Produkt folgt.

Die Division ist die Umkehrung der Multiplikation. Multiplizieren wir den Ausdruck $\frac{r_1}{r_2} (\cos (\varphi_1 - \varphi_2) + i \cdot \sin (\varphi_1 - \varphi_2))$ mit $r_2 (\cos \varphi_2 + i \cdot \sin \varphi_2)$, so erhalten wir nach der eben bewiesenen Formel

$$\frac{r_1}{r_2} \cdot r_2 (\cos (\varphi_1 - \varphi_2 + \varphi_2) + i \cdot \sin (\varphi_1 - \varphi_2 + \varphi_2)) = r_1 (\cos \varphi_1 + i \cdot \sin \varphi_1).$$

Da der Quotient zweier komplexer Zahlen eindeutig bestimmt ist, ist die Richtigkeit der zweiten Formel ebenfalls gezeigt.

Die Multiplikation und die Division komplexer Zahlen in trigonometrischer Form lassen sich in der Gaußschen Zahlenebene veranschaulichen (vgl. Abb. 48 und Abb. 49).

Der dem Produkt zugeordnete Punkt liegt auf dem Strahl, dessen Argument die Summe der Winkel φ_1 und φ_2 ist, und hat den Betrag $r_1 \cdot r_2$ als Abstand vom

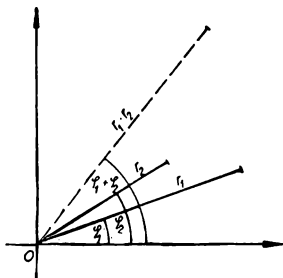


Abb. 48

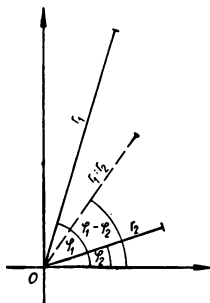


Abb. 49

Nullpunkt. Da nach dem Strahlensatz die entsprechenden Seiten ähnlicher Dreiecke einander proportional sind, gilt

$$(r_1 \cdot r_2) : r_1 = r_2 : 1;$$

also kann die Länge $r_1 \cdot r_2$ konstruiert werden (vgl. Abb. 50). Wir verbinden den der Zahl $r_1 (\cos \varphi_1 + i \cdot \sin \varphi_1)$ zugeordneten Punkt Z_1 mit dem die 1 darstellenden Punkt E und erhalten das Dreieck OEZ_1 . Der Punkt Z_2 ist das Bild der Zahl $r_2 (\cos \varphi_2 + i \cdot \sin \varphi_2)$. Die freien Schenkel der in O und Z_2 an $\overline{OZ_2}$ angelegten Winkel φ_1 bzw. $\angle OEZ_1$ schneiden einander in Z . Es ist $\triangle OEZ_1 \sim \triangle OZ_2Z$ und damit die Länge von $\overline{OZ}$ gleich dem Produkt $r_1 \cdot r_2$. Der konstruierte Punkt Z stellt also das Produkt

$$r_1 r_2 (\cos (\varphi_1 + \varphi_2) + i \cdot \sin (\varphi_1 + \varphi_2))$$

dar.

Entsprechend liegt der dem Quotienten zweier komplexer Zahlen zugeordnete Punkt auf dem Strahl, dessen Argument die Differenz der Winkel φ_1 und φ_2 ist.

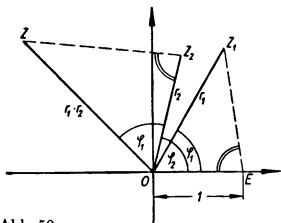


Abb. 50

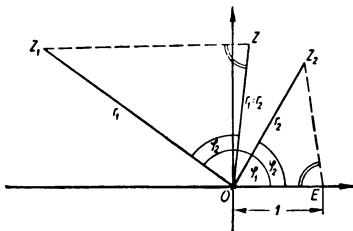


Abb. 51

Er hat den Betrag $\frac{r_1}{r_2}$ als Abstand vom Nullpunkt. Wir können ihn, den eben erläuterten Gedankengang umkehrend, ebenfalls konstruieren (vgl. Abb. 51).

Es ist dann $\triangle OZ_2E \sim \triangle OZ_1Z$ und wegen

$$\frac{r_1}{r_2} : 1 = r_1 : r_2$$

die Länge von $\overline{OZ}$ der Quotient $\frac{r_1}{r_2}$. Der Punkt Z stellt den Quotienten

$$\frac{r_1}{r_2} (\cos(\varphi_1 - \varphi_2) + i \cdot \sin(\varphi_1 - \varphi_2))$$

dar.

Aufgaben

1. a) $2 \left(\cos \frac{\pi}{6} + i \cdot \sin \frac{\pi}{6} \right) \cdot \frac{2}{3} \left(\cos \frac{\pi}{3} + i \cdot \sin \frac{\pi}{3} \right)$
 b) $\frac{1}{2} (\cos \pi + i \cdot \sin \pi) \cdot \frac{4}{5} \left(\cos \frac{3\pi}{4} + i \cdot \sin \frac{3\pi}{4} \right)$
 c) $7,5 (\cos 32^\circ + i \cdot \sin 32^\circ) \cdot 4,8 (\cos 117^\circ + i \cdot \sin 117^\circ)$
 d) $5 \left(\cos \frac{7\pi}{8} + i \cdot \sin \frac{7\pi}{8} \right) \cdot 3,5 \left(\cos \frac{2\pi}{3} + i \cdot \sin \frac{2\pi}{3} \right)$
 e) $10 (\cos 53,1^\circ + i \cdot \sin 53,1^\circ) \cdot 12 (\cos 342,7^\circ + i \cdot \sin 342,7^\circ)$
2. a) $\frac{2}{3} \left(\cos \frac{\pi}{3} + i \cdot \sin \frac{\pi}{3} \right) : 2 \left(\cos \frac{\pi}{6} + i \cdot \sin \frac{\pi}{6} \right)$
 b) $2 \left(\cos \frac{\pi}{6} + i \cdot \sin \frac{\pi}{6} \right) : \frac{2}{3} \left(\cos \frac{\pi}{3} + i \cdot \sin \frac{\pi}{3} \right)$
 c) $4 (\cos 184^\circ + i \cdot \sin 184^\circ) : 12 (\cos 254^\circ + i \cdot \sin 254^\circ)$
 d) $1 : \frac{4}{3} \left(\cos \frac{7\pi}{3} + i \cdot \sin \frac{7\pi}{3} \right)$
 e) $4i : 5 (\cos 45^\circ + i \cdot \sin 45^\circ)$
 f) $-8 : \sqrt{8} \left(\cos \frac{5\pi}{4} + i \cdot \sin \frac{5\pi}{4} \right)$

8. Führen Sie einige der in 1 und 2 gestellten Aufgaben zeichnerisch durch!

12. Potenzen und Wurzeln komplexer Zahlen, der Satz von Moivre

Die Multiplikation komplexer Zahlen in trigonometrischer Form ergibt eine übersichtliche Schreibweise für die Potenzen komplexer Zahlen. So ist zum Beispiel

$$\begin{aligned} \{r(\cos \varphi + i \cdot \sin \varphi)\}^2 &= \{r(\cos \varphi + i \cdot \sin \varphi)\} \cdot \{r(\cos \varphi + i \cdot \sin \varphi)\} \\ &= r^2 (\cos 2\varphi + i \cdot \sin 2\varphi) \end{aligned}$$

$$\begin{aligned} \{r(\cos \varphi + i \cdot \sin \varphi)\}^3 &= \{r(\cos \varphi + i \cdot \sin \varphi)\}^2 \cdot \{r(\cos \varphi + i \cdot \sin \varphi)\} \\ &= r^3 (\cos 3\varphi + i \cdot \sin 3\varphi). \end{aligned}$$

Wir wollen sogleich die nach **Moivre** benannte Formel

$$\{r(\cos \varphi + i \cdot \sin \varphi)\}^n = r^n(\cos n\varphi + i \cdot \sin n\varphi)$$

für positive, ganzzahlige n beweisen. Wir führen einen Induktionsschluß durch und nehmen an, daß die Formel bereits für n bewiesen sei. Sie gilt dann wegen

$$\begin{aligned}\{r(\cos \varphi + i \cdot \sin \varphi)\}^{n+1} &= \{r(\cos \varphi + i \cdot \sin \varphi)\}^n \cdot \{r(\cos \varphi + i \cdot \sin \varphi)\} \\ &= \{r^n(\cos n\varphi + i \cdot \sin n\varphi)\} \cdot \{r(\cos \varphi + i \cdot \sin \varphi)\} \\ &= r^n r(\cos(n\varphi + \varphi) + i \cdot \sin(n\varphi + \varphi)) \\ &= r^{n+1}(\cos(n+1)\varphi + i \cdot \sin(n+1)\varphi)\end{aligned}$$

auch für $n+1$. Für $n=1$ ist sie sicher richtig. Sie gilt also für alle positiven ganzzahligen n .

Es ist dabei wieder zu beachten, daß das Argument $n\varphi$ größer als 2π werden kann und dann um ganzzahlige Vielfache der Periode 2π abzuändern ist.

Satz von Moivre:

Die Formel

$$\{r(\cos \varphi + i \cdot \sin \varphi)\}^n = r^n(\cos n\varphi + i \cdot \sin n\varphi)$$

gilt für beliebige rationale Werte von n .

Beweis:

Für positive ganzzahlige n haben wir die Formel bereits als richtig erkannt. Im folgenden werden wir sie in drei Teilschritten für alle anderen rationalen Werte von n beweisen.

a) Beweis der Moivreschen Formel für negative ganzzahlige $n = -p$

Nach der Definition von Potenzen mit negativen Exponenten ist

$$\{r(\cos \varphi + i \cdot \sin \varphi)\}^{-p} = \frac{1}{\{r(\cos \varphi + i \cdot \sin \varphi)\}^p}.$$

Wir erweitern mit $(\cos \varphi - i \cdot \sin \varphi)^p$ und erhalten:

$$\{r(\cos \varphi + i \cdot \sin \varphi)\}^{-p} = \frac{(\cos \varphi - i \cdot \sin \varphi)^p}{r^p \{(\cos \varphi + i \cdot \sin \varphi)(\cos \varphi - i \cdot \sin \varphi)\}^p}.$$

Da $\cos^2 \varphi + \sin^2 \varphi = 1$ ist, steht im Nenner nur r^p . Den Zähler können wir aber wegen $\cos \varphi = \cos(-\varphi)$ und $-\sin \varphi = \sin(-\varphi)$ in $(\cos(-\varphi) + i \cdot \sin(-\varphi))^p = \cos(-p\varphi) + i \cdot \sin(-p\varphi)$ umwandeln.

Es ist also

$$\{r(\cos \varphi + i \cdot \sin \varphi)\}^{-p} = r^{-p}(\cos(-p\varphi) + i \cdot \sin(-p\varphi)).$$

Damit ist die Gültigkeit der Formel für negative ganzzahlige Exponenten bewiesen.

b) Beweis der Moivreschen Formel für positive Brüche der Form $n = \frac{1}{q}$

Wir bilden von der komplexen Zahl $r^{\frac{1}{q}} \left(\cos \frac{1}{q} \varphi + i \cdot \sin \frac{1}{q} \varphi \right)$ die q -te Potenz. Es ist

$$\left\{ r^{\frac{1}{q}} \left(\cos \frac{1}{q} \varphi + i \cdot \sin \frac{1}{q} \varphi \right) \right\}^q = \left(r^{\frac{1}{q}} \right)^q \left(\cos \frac{q}{q} \varphi + i \cdot \sin \frac{q}{q} \varphi \right) \\ = r (\cos \varphi + i \cdot \sin \varphi).$$

Erheben wir diese Gleichung in die $\frac{1}{q}$ -te Potenz, so erhalten wir mit vertauschten Seiten

$$\{r(\cos \varphi + i \cdot \sin \varphi)\}^{\frac{1}{q}} = r^{\frac{1}{q}} \left(\cos \frac{1}{q} \varphi + i \cdot \sin \frac{1}{q} \varphi \right).$$

Damit ist die Gültigkeit der Formel für Exponenten von der Form $n = \frac{1}{q}$ bewiesen.

c) Beweis der Moivreschen Formel für beliebige rationale n .

Da jede rationale Zahl $n = \frac{p}{q}$ als Quotient zweier ganzer Zahlen mit $q > 0$ geschrieben werden kann, folgt aus den letzten beiden Teilschritten bereits die Gültigkeit der Moivreschen Formel für beliebige rationale n :

$$\{r(\cos \varphi + i \cdot \sin \varphi)\}^{\frac{p}{q}} = \left[\{r(\cos \varphi + i \cdot \sin \varphi)\}^{\frac{1}{q}} \right]^p = \left[r^{\frac{1}{q}} \left(\cos \frac{1}{q} \varphi + i \cdot \sin \frac{1}{q} \varphi \right) \right]^p \\ = r^{\frac{p}{q}} \left(\cos \frac{p}{q} \varphi + i \cdot \sin \frac{p}{q} \varphi \right).$$

Damit ist der Satz von Moivre bewiesen.

Es lassen sich also mit Hilfe der Moivreschen Formel alle Potenzen und Wurzeln komplexer Zahlen mit rationalen Exponenten berechnen. Wir können hier nur angeben, daß auch alle Wurzeln und Potenzen mit reellem, ja sogar mit komplexem Exponenten n innerhalb des Bereiches der komplexen Zahlen existieren. Auch für sie gilt die Moivresche Formel.

Aufgaben

1. Berechnen Sie die folgenden Potenzen komplexer Zahlen einmal in der gegebenen arithmetischen Form, sodann in trigonometrischer Form nach dem Moivreschen Satz!

a) $(4 + 3i)^3$

b) $\left(\frac{1}{2} + \frac{1}{2} i \sqrt{3} \right)^3$

c) $(2 + 2i)^4$

2. Verfolgen Sie die Rechnungen der Aufgabe 1 in der Gaußschen Zahlenebene! Sprechen Sie das Ergebnis in allgemeiner Form für beliebige Potenzen aus!

3. a) $\left\{ \frac{1}{2} (\cos \pi + i \cdot \sin \pi) \right\}^3$ b) $\left\{ \frac{2}{3} \left(\cos \frac{\pi}{3} + i \cdot \sin \frac{\pi}{3} \right) \right\}^7$
 c) $\left\{ 2 \left(\cos \frac{\pi}{6} + i \cdot \sin \frac{\pi}{6} \right) \right\}^5$ d) $\left\{ 1 \left(\cos \frac{\pi}{3} + i \cdot \sin \frac{\pi}{3} \right) \right\}^6$
 e) $\{ 5 (\cos 10^\circ + i \cdot \sin 10^\circ) \}^4$ f) $\{ 2 (\cos 127^\circ + i \cdot \sin 127^\circ) \}^8$

4. Berechnen Sie die folgenden Wurzeln nach dem Moivreschen Satz für gebrochene Exponenten und prüfen Sie die Resultate! Die Mehrdeutigkeit dieser Aufgaben wird erst im folgenden Abschnitt behandelt und soll vorläufig außer acht gelassen werden.

- a) $\sqrt[2]{9 \left(\cos \frac{\pi}{2} + i \cdot \sin \frac{\pi}{2} \right)}$ b) $\sqrt[4]{32 \left(\cos \frac{\pi}{3} + i \cdot \sin \frac{\pi}{3} \right)}$
 c) $\sqrt[5]{32 \left(\cos \frac{\pi}{2} + i \cdot \sin \frac{\pi}{2} \right)}$ d) $\sqrt[3]{27 \left(\cos \frac{3}{2}\pi + i \cdot \sin \frac{3}{2}\pi \right)}$
 e) $\sqrt[12]{64 (\cos 144^\circ + i \cdot \sin 144^\circ)}$ f) $\sqrt[3]{4 (\cos 260^\circ + i \cdot \sin 260^\circ)}$

13. Algebraische Gleichungen

Eine algebraische Gleichung n -ten Grades hat die Form

$$a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0 = 0.$$

Die Koeffizienten $a_n, a_{n-1}, \dots, a_0$ können beliebige reelle oder auch komplexe Zahlen sein, wobei nur die Einschränkung $a_n \neq 0$ notwendig ist. Für die Theorie der algebraischen Gleichungen sind die komplexen Zahlen ein unentbehrliches Hilfsmittel.

So können schon bei quadratischen Gleichungen mit reellen Koeffizienten komplexe Wurzeln auftreten. Im Lehrbuch der Mathematik, 9. Schuljahr, wurden für die quadratischen Gleichungen in Normalform

$$x^2 + px + q = 0$$

die Wurzeln

$$x_1 = -\frac{p}{2} + \sqrt{\frac{p^2}{4} - q},$$

$$x_2 = -\frac{p}{2} - \sqrt{\frac{p^2}{4} - q}$$

abgeleitet. Der Radikand der Quadratwurzel, die Diskriminante $D = \frac{p^2}{4} - q$, gibt an, welcher Art diese Wurzeln sind. So sind für $D > 0$ die Wurzeln x_1 und x_2 reell und voneinander verschieden, für $D = 0$ fallen beide zu einer reellen Doppel-

wurzel zusammen. Im Falle $D < 0$ konnten wir bisher nur feststellen, daß **keine** reellen Wurzeln existieren. Es wird nämlich $\sqrt{D}$ imaginär, und wir erhalten die beiden Lösungen

$$x_1 = -\frac{p}{2} + i \sqrt{\left|\frac{p^2}{4} - q\right|},$$

$$x_2 = -\frac{p}{2} - i \sqrt{\left|\frac{p^2}{4} - q\right|}.$$

Sie sind konjugiert komplex.

Wir können also feststellen, daß im Bereich der komplexen Zahlen eine quadratische Gleichung stets zwei Wurzeln hat, die allerdings zu einer Doppelwurzel zusammenfallen können. Dieser Satz bleibt richtig, wenn man quadratische Gleichungen mit komplexen Koeffizienten betrachtet.

Für die Gleichungen dritten Grades gibt es ähnliche, jedoch schon wesentlich kompliziertere Auflösungsformeln. Sie werden nach dem italienischen Mathematiker Cardano benannt und sagen aus, daß jede Gleichung dritten Grades genau drei Wurzeln hat, von denen allerdings wieder zwei zu einer doppelten oder alle drei zu einer dreifachen Wurzel zusammenfallen können. Wir wollen dies an einigen Beispielen zeigen.

So hat die Gleichung

$$x^3 - 6x^2 + 11x - 6 = (x-1)(x-2)(x-3) = 0$$

die drei voneinander verschiedenen Wurzeln

$$x_1 = 1, \quad x_2 = 2, \quad x_3 = 3$$

(vgl. Lehrbuch der Mathematik, 11. Schuljahr, S. 38), die Gleichung

$$x^3 - 2x^2 + 2x = (x-0)(x-(1+i))(x-(1-i)) = 0$$

die reelle Wurzel $x_1 = 0$ und die beiden konjugiert komplexen Wurzeln $x_2 = 1 + i$, $x_3 = 1 - i$, die Gleichung

$$x^3 - 3x^2 + 4 = (x+1)(x-2)(x-2) = 0$$

die reelle Wurzel $x_1 = -1$ und die reelle Doppelwurzel $x_2 = x_3 = 2$ und schließlich die Gleichung

$$x^3 - 3x^2 + 3x - 1 = (x-1)(x-1)(x-1) = 0$$

eine reelle dreifache Wurzel $x_1 = x_2 = x_3 = 1$.

Auch bei Gleichungen dritten Grades ist also der Satz, daß es drei Wurzeln gibt, nur richtig, wenn man die komplexen Wurzeln mitberücksichtigt. Sie müssen übrigens bei einer Gleichung dritten Grades mit reellen Koeffizienten konjugiert komplex auftreten, so daß eine solche Gleichung entweder drei reelle oder eine reelle und zwei konjugiert komplexe Wurzeln hat.

Eine Verallgemeinerung dieser Behauptung ist der berühmte **Fundamentalsatz der Algebra**, der zum ersten Male von Carl Friedrich Gauß im Jahre 1799 bewiesen wurde:

Jede algebraische Gleichung n -ten Grades mit reellen oder komplexen Koeffizienten hat im Bereich der komplexen Zahlen genau n Wurzeln, wenn man jede Wurzel mit ihrer Vielfachheit zählt.

Man kann also jede algebraische Gleichung n -ten Grades in n Wurzelfaktoren zerlegen:

$$a_n x^n + \dots + a_1 x^1 + a_0 = a_n (x - x_1) (x - x_2) (x - x_3) \dots (x - x_n) = 0.$$

Dabei sind $x_1, \dots, x_n$ die Wurzeln dieser Gleichung. Sie können reell oder komplex sein und können mehrfach auftreten. Den Fundamentalsatz der Algebra können wir mit den zur Verfügung stehenden Mitteln nicht beweisen. Wir wollen ihn jedoch für einige Gleichungen nachprüfen.

So muß die Gleichung n -ten Grades

$$x^n - 1 = 0$$

n Wurzeln haben, die man wegen

$$x = \sqrt[n]{1}$$

die n -ten **Einheitswurzeln** nennt. Im Bereich der reellen Zahlen hat diese Gleichung aber nur die Wurzel $+1$, falls n ungerade, und die Wurzeln $+1$ und -1 , falls n gerade ist. Wir wollen die komplexen Wurzeln der Gleichung ermitteln.

Dazu schreiben wir 1 als komplexe Zahl in trigonometrischer Form und wenden den Moivreschen Satz an:

$$1 = r (\cos \varphi + i \cdot \sin \varphi),$$

$$\sqrt[n]{1} = 1^{\frac{1}{n}} = r^{\frac{1}{n}} \left(\cos \frac{\varphi}{n} + i \cdot \sin \frac{\varphi}{n} \right).$$

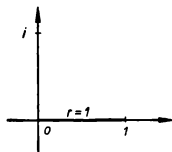


Abb. 52

Dabei ist $r = 1$. Da $r^{\frac{1}{n}}$ als Betrag einer komplexen Zahl positiv reell werden muß, ist $r^{\frac{1}{n}} = \sqrt[n]{1} = +1$ zu wählen; das Argument φ ist zunächst 0 (vgl. Abb. 52). Dann wird aber $\frac{\varphi}{n}$ ebenfalls 0 , und wir erhalten

$$\sqrt[n]{1} = 1 \cdot (\cos 0 + i \cdot \sin 0) = (1 + 0 \cdot i) = 1.$$

Nun benutzen wir die Periodizität der Sinus- und Kosinusfunktion. Wegen $\sin 0 = \sin 2\pi$ und $\cos 0 = \cos 2\pi$ können wir auch 2π als Argument φ für die Zahl 1 wählen: Dann wird $\frac{\varphi}{n} = \frac{2\pi}{n}$ und damit

$$\sqrt[n]{1} = 1 \cdot \left(\cos \frac{2\pi}{n} + i \cdot \sin \frac{2\pi}{n} \right).$$

Diese komplexe Zahl liegt auf dem Kreis mit dem Radius 1 und dem Nullpunkt der Gaußschen Zahlenebene als Mittelpunkt, um $\frac{1}{n}$ des Vollkreises im mathematisch positiven Sinne aus der reellen Achse herausgedreht. Sie erfüllt tatsächlich die Gleichung

$$x^n = 1.$$

Das veranlaßt uns, das Argument φ von 1 der Reihe nach als

$$\varphi = 2 \cdot 2\pi, \quad \varphi = 3 \cdot 2\pi, \quad \varphi = (n-1) \cdot 2\pi$$

anzusetzen. Damit erhält man für $\sqrt[n]{1}$ die weiteren Argumente

$$\frac{\varphi}{n} = 2 \cdot \frac{2\pi}{n}, \quad \frac{\varphi}{n} = 3 \cdot \frac{2\pi}{n}, \quad \frac{\varphi}{n} = (n-1) \cdot \frac{2\pi}{n}$$

und damit insgesamt n verschiedene Lösungen

$$\sqrt[n]{1} = 1 \cdot \left(\cos k \frac{2\pi}{n} + i \cdot \sin k \frac{2\pi}{n} \right) \quad (k = 0, 1, \dots, n-1)$$

unserer Gleichung. Für alle anderen ganzzahligen k ergeben sich keine neuen Werte für die Wurzeln. In der Abbildung 53 sind für $n = 6$ und in der Abbildung 54 für $n = 7$ die n -ten Einheitswurzeln in der komplexen Zahlenebene dargestellt.

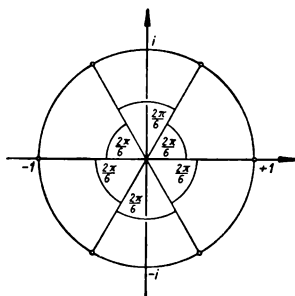


Abb. 53

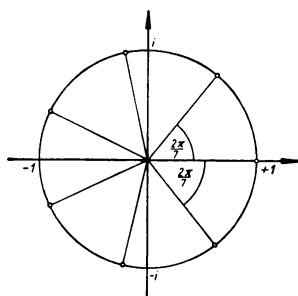


Abb. 54

Im allgemeinen Falle sind es die Eckpunkte eines in den Einheitskreis einbeschriebenen regulären n -Eckes, dessen eine Ecke im Punkte $+1$ liegt. Es bestätigt sich, daß -1 nur für gerade n Lösung ist. Die nicht reellen Wurzeln treten konjugiert komplex auf. Dieser geometrischen Bedeutung wegen nennen wir die Gleichung $x^n - 1 = 0$ auch **Kreisteilungsgleichung**.

Die gleichen Überlegungen können wir auch für die Gleichung

$$x^n - a = 0$$

mit beliebigem, positiv reellem a durchführen. Die $\sqrt[n]{a}$ muß also ebenfalls n Lösungen haben. Wir erhalten sie, wenn wir a als komplexe Zahl in trigonometrischer Form schreiben

$$a = r (\cos \varphi + i \cdot \sin \varphi) = a (\cos \varphi + i \cdot \sin \varphi)$$

und bei Anwendung des Moivreschen Satzes das Argument φ der Reihe nach als $\varphi = 0, \varphi = 2\pi, \varphi = 2 \cdot 2\pi, \dots, \varphi = (n-1) \cdot 2\pi$ ansetzen. Es sind dann

$$\sqrt[n]{a} = a^{\frac{1}{n}} \left(\cos k \frac{2\pi}{n} + i \cdot \sin k \frac{2\pi}{n} \right) \quad (k = 0, 1, \dots, n-1)$$

die n verschiedenen Lösungen unserer Gleichung, wobei für $a^{\frac{1}{n}}$ als Betrag einer komplexen Zahl die eindeutig bestimmte, positive reelle n -te Wurzel aus a zu setzen ist.

So hat zum Beispiel die $\sqrt[3]{8}$ die drei Lösungen

$$\sqrt[3]{8} = 8^{\frac{1}{3}} \left(\cos 0 \cdot \frac{2\pi}{3} + i \cdot \sin 0 \cdot \frac{2\pi}{3} \right) = 2 (1 + 0) = 2,$$

$$\sqrt[3]{8} = 8^{\frac{1}{3}} \left(\cos 1 \cdot \frac{2\pi}{3} + i \cdot \sin 1 \cdot \frac{2\pi}{3} \right) = 2 \left(\cos \frac{2\pi}{3} + i \cdot \sin \frac{2\pi}{3} \right),$$

$$\sqrt[3]{8} = 8^{\frac{1}{3}} \left(\cos 2 \cdot \frac{2\pi}{3} + i \cdot \sin 2 \cdot \frac{2\pi}{3} \right) = 2 \left(\cos \frac{4\pi}{3} + i \cdot \sin \frac{4\pi}{3} \right),$$

die in der Abbildung 55 veranschaulicht sind.

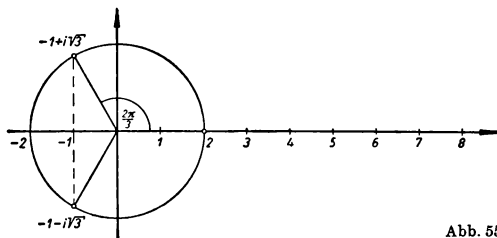


Abb. 55

Vergleichen wir das Ergebnis für $\sqrt[n]{a}$ mit dem für $\sqrt[n]{1}$, so sehen wir, daß wir alle Lösungen $\sqrt[n]{a}$ erhalten können, indem wir die positive reelle n -te Wurzel aus a mit allen n -ten Einheitswurzeln multiplizieren. Wir können sogar von irgendeiner der

n Lösungen von $\sqrt[n]{a}$ ausgehen und durch Multiplikation dieser Lösung mit allen n -ten Einheitswurzeln alle anderen erhalten. Ist nämlich α eine beliebige Lösung der $\sqrt[n]{a}$, das heißt $\alpha^n = a$, so erfüllt auch stets $\alpha \cdot \sqrt[n]{1}$ wegen

$$(\alpha \cdot \sqrt[n]{1})^n = \alpha^n (\sqrt[n]{1})^n = a \cdot 1 = a$$

die Gleichung $x^n = a$.

Wir wollen schließlich zeigen, daß jede komplexe Zahl

$$r (\cos \varphi + i \cdot \sin \varphi)$$

genau n voneinander verschiedene n -te Wurzeln hat. Es sind dies die n komplexen Zahlen

$$r^{\frac{1}{n}} \left(\cos \frac{\varphi + k \cdot 2\pi}{n} + i \cdot \sin \frac{\varphi + k \cdot 2\pi}{n} \right) \quad (k = 0, 1, \dots, n-1),$$

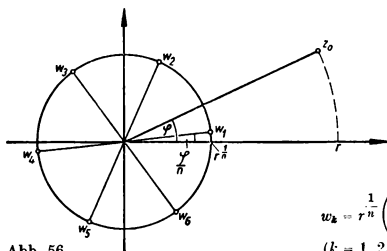


Abb. 56

Zur Abbildung:

Es ist $z_0 = r (\cos \varphi + i \sin \varphi)$ und

$$w_k = r^{\frac{1}{n}} \left(\cos \frac{\varphi + (k-1) 2\pi}{n} + i \sin \frac{\varphi + (k-1) 2\pi}{n} \right), \quad (k = 1, 2, \dots, 6).$$

deren Betrag $r^{\frac{1}{n}}$ die positive reelle n -te Wurzel des Betrages r ist. Wir erhalten sie, indem wir entweder vor Anwendung des Moivreschen Lehrsatzes das Argument φ der Reihe nach als

$$\varphi + 0 \cdot 2\pi, \varphi + 1 \cdot 2\pi, \dots, \varphi + (n-1) \cdot 2\pi$$

ansetzen oder die sich aus $r (\cos \varphi + i \cdot \sin \varphi)$ nach dem Moivreschen Satz ergebende Lösung

$$r^{\frac{1}{n}} \left(\cos \frac{\varphi}{n} + i \cdot \sin \frac{\varphi}{n} \right)$$

mit allen n -ten Einheitswurzeln multiplizieren:

$$\left\{ r^{\frac{1}{n}} \left(\cos \frac{\varphi}{n} + i \cdot \sin \frac{\varphi}{n} \right) \right\} \cdot \left\{ \left(\cos k \frac{2\pi}{n} + i \cdot \sin k \frac{2\pi}{n} \right) \right\} \\ = r^{\frac{1}{n}} \left(\cos \frac{\varphi + k \cdot 2\pi}{n} + i \cdot \sin \frac{\varphi + k \cdot 2\pi}{n} \right). \quad (k = 0, 1, \dots, n-1)$$

In der Abbildung 56 sind als Beispiel die sechsten Wurzeln der komplexen Zahl $r (\cos \varphi + i \cdot \sin \varphi)$ dargestellt.

Aufgaben

1. Lösen Sie die folgenden quadratischen Gleichungen!

a) $x^2 + x + 1 = 0$

b) $8x^2 - 8x + 10 = 0$

c) $4x^2 + 12x + 10 = 0$

d) $4x^2 + 12x + 9 = 0$

e) $4x^2 + 12x + 8 = 0$

f) $x^2 - 5x + 2,25 = 0$

2. Die im Text angegebenen Gleichungen dritten Grades sind nachzuprüfen.

3. Die Gleichung vierten Grades $x^4 + 2x^3 - x - 2 = 0$ hat zwei reelle und zwei konjugiert komplexe Wurzeln. Sie sind zu bestimmen.

Anleitung: Die reellen Wurzeln x_1 und x_2 sind durch Probieren leicht zu finden, die Division durch $(x - x_1)(x - x_2)$ führt zu einer Gleichung zweiten Grades.

4. Berechnen Sie die dritten, vierten, fünften und sechsten Einheitswurzeln!

5. Sämtliche Lösungen von

$$\sqrt[3]{625}, \sqrt[3]{-8}, \sqrt[3]{27}, \sqrt[3]{-27}, \sqrt[3]{2}, \sqrt[3]{-32}$$

sind als komplexe Zahlen und ihre Bilder in der Gaußschen Zahlenebene anzugeben.

6. Die Lösungen der Aufgabe 4 zu Abschnitt 12 sind vollständig zu berechnen.

14. Anwendung komplexer Zahlen in Physik und Technik

Die komplexen Zahlen sind in den Anwendungsgebieten der Mathematik von großer Bedeutung. Es gibt heute wenige Gebiete der Physik und der Technik, die auf das Hilfsmittel der komplexen Zahlen verzichten könnten. Wir müssen uns darauf beschränken, an einigen Beispielen die Anwendbarkeit der komplexen Zahlen auf physikalische Probleme zu zeigen.

a) Der Wechselstromwiderstand

Es ist bekannt, daß in einer von Wechselstrom durchflossenen Spule außer dem Ohmschen, auch für Gleichstrom vorhandenen Widerstand R_Ω der durch Induktion hervorgerufene Blindwiderstand (induktiver Widerstand) $R_L = 2\pi f \cdot L$ auftritt. Er hängt von der Selbstinduktion L der Spule und der Frequenz f des Wechselstromes ab. Wir können den Wechselstromwiderstand als eine komplexe Zahl $R = R_\Omega + iR_L$ auffassen (vgl. Abb. 57).

Der Betrag $|R| = \sqrt{R_\Omega^2 + R_L^2}$ ist der in das Ohmsche Gesetz eingehende Scheinwiderstand, das Argument $\varphi = \arctg \frac{R_L}{R_\Omega}$

die Phasenverschiebung. Der Vorteil

dieser Betrachtungsweise ist, daß sich bei Hintereinanderschaltung mehrerer Spulen die Widerstände wie komplexe Zahlen addieren und bei Parallelschaltung

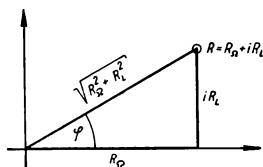


Abb. 57

die Kirchhoffsche Regel gilt, wenn wir für alle auftretenden Widerstände $R_{\Omega} + i R_L$ einsetzen.

Wir wollen dazu zwei Beispiele durchrechnen:

1. An den Stromkreis (vgl. Abb. 58) sei Netzwechselspannung ($f = 50$ Hz, $U = 220$ V) angelegt. Die Spule S_1 habe eine Selbstinduktion $L_1 = 7$ Henry, die Spule S_2 eine Selbstinduktion $L_2 = 3,5$ Henry und beide einen Ohmschen Widerstand $R_{\Omega_1} = R_{\Omega_2} = 200 \Omega$. Wir wollen die Stromstärke berechnen. Es ist

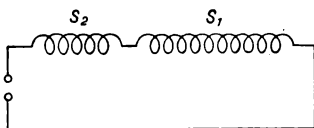


Abb. 58

$$R_1 = R_{\Omega_1} + i 2\pi f L_1 = 200 + i 2\pi \cdot 50 \cdot 7 \approx 200 + 2200 i$$

$$R_2 = R_{\Omega_2} + i 2\pi f L_2 = 200 + i 2\pi \cdot 50 \cdot 3,5 \approx 200 + 1100 i,$$

wobei wir zur Vereinfachung für π den Näherungswert $\frac{22}{7}$ gesetzt haben. Dann ist $R = R_1 + R_2 \approx 400 + 3300 i$ der komplexe Gesamtwiderstand des Kreises. Der Betrag der komplexen Zahl R ist $\sqrt{400^2 + 3300^2} \approx 3324$. Damit ist die Stromstärke

$$I = \frac{U}{|R|} \approx \frac{220}{3324} \text{ A} \approx 0,066 \text{ A}.$$

2. Schalten wir dagegen beide Spulen parallel (vgl. Abb. 59), so haben wir nach der Kirchhoffschen Regel

$$R = \frac{R_1 R_2}{R_1 + R_2} \approx \frac{(200 + 2200 i)(200 + 1100 i)}{400 + 3300 i}.$$

Wir bringen die drei komplexen Zahlen in ihre trigonometrische Form und wenden den Satz von Moivre an. Da wir lediglich den Betrag von R brauchen, verzichten wir auf die Berechnung der Ausdrücke $\cos \varphi + i \sin \varphi$. Es ist dann

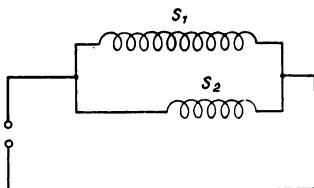


Abb. 59

$$|R| = \frac{|R_1| \cdot |R_2|}{|R_1 + R_2|} \approx \frac{2210 \cdot 1120}{3324} \approx 745,$$

$$I = \frac{U}{|R|} \approx \frac{220}{745} \text{ A} \approx 0,295 \text{ A}.$$

Auf ähnliche Weise verfährt man, wenn ein Kondensator in einem Wechselstromkreis liegt.

b) Bewegung auf einem Kreis

Auf einem Kreis mit dem Radius r bewege sich ein Körper mit der konstanten Geschwindigkeit v (vgl. Abb. 60). Wir können dann seinen Ort s zu jedem Zeitpunkt t durch die komplexe Zahl $s = r \left(\cos \frac{v}{r} t + i \cdot \sin \frac{v}{r} t \right)$ kennzeichnen, wobei der Winkel φ im Bogenmaß das Verhältnis des zurückgelegten Bogens $v \cdot t$ zum Kreisradius r ist.

Da die Geschwindigkeit die erste, die Beschleunigung die zweite Ableitung des Weges nach der Zeit ist, und, wie wir nicht beweisen wollen, die Differentiation von Funktionen im Bereich der komplexen Zahlen ebenso wie im Bereich der reellen Zahlen verläuft, können wir auch hier sofort die Geschwindigkeit und die Zentralbeschleunigung berechnen. Es ist dann

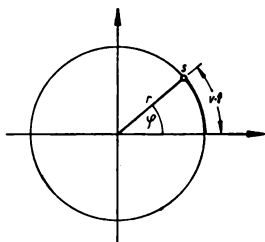


Abb. 60

$$\frac{ds}{dt} = r \left(-\sin \frac{v}{r} t + i \cdot \cos \frac{v}{r} t \right) \cdot \frac{v}{r},$$

$$\frac{d^2s}{dt^2} = r \left(-\cos \frac{v}{r} t - i \cdot \sin \frac{v}{r} t \right) \cdot \frac{v^2}{r^2}.$$

Berücksichtigen wir in beiden komplexen Zahlen nur die Beträge, so erhalten wir tatsächlich die Geschwindigkeit v aus der ersten, die Zentralbeschleunigung $\frac{v^2}{r}$ aus der zweiten Ableitung.

c) Schwingungen eines Massenpunktes

Eine punktförmige Masse m schwinde zwischen zwei gleich starken Federn (vgl. Abb. 61).

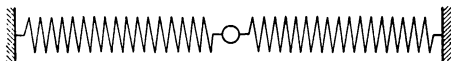


Abb. 61

Sehen wir von Reibungsverlusten ab, so ist die rücktreibende Kraft K um so größer, je mehr der Punkt aus der Mittellage ausgelenkt ist. Dabei ist K proportional der Entfernung x vom Ruhepunkt. Unter Verwendung eines Proportionalitätsfaktors $- \kappa$ ergibt sich demnach

$$K = - \kappa \cdot x.$$

Da die Kraft das Produkt aus Masse und Beschleunigung ist, erhalten wir eine Relation

$$-\kappa \cdot x = m \cdot \frac{d^2 x}{dt^2}.$$

Ähnliche Beziehungen treten bei allen Schwingungsproblemen auf. Die eben gewonnene Relation wird von der komplexen Zahl

$$x = r \left(\cos \sqrt{\frac{\kappa}{m}} \cdot t + i \cdot \sin \sqrt{\frac{\kappa}{m}} \cdot t \right)$$

als Funktion von t erfüllt. Es ist nämlich

$$\frac{dx}{dt} = r \left(-\sin \sqrt{\frac{\kappa}{m}} \cdot t + i \cdot \cos \sqrt{\frac{\kappa}{m}} \cdot t \right) \cdot \sqrt{\frac{\kappa}{m}},$$

$$\frac{d^2 x}{dt^2} = r \left(-\cos \sqrt{\frac{\kappa}{m}} \cdot t - i \cdot \sin \sqrt{\frac{\kappa}{m}} \cdot t \right) \left(\sqrt{\frac{\kappa}{m}} \right)^2,$$

$$\frac{d^2 x}{dt^2} = -\frac{\kappa}{m} r \left(\cos \sqrt{\frac{\kappa}{m}} \cdot t + i \cdot \sin \sqrt{\frac{\kappa}{m}} \cdot t \right),$$

das heißt

$$\frac{d^2 x}{dt^2} = -\frac{\kappa}{m} \cdot x.$$

Dabei entspricht der komplexen Zahl x ein Punkt in der Gaußschen Zahlenebene, der auf dem Kreis vom Radius r mit der Geschwindigkeit $r \cdot \sqrt{\frac{\kappa}{m}}$ umläuft (vgl. Abb. 62).

Die Projektion dieses Punktes auf die reelle Achse, also der Realteil von x , beschreibt den Weg der schwingenden Masse. Der Betrag r ist die Amplitude, das Argument $\sqrt{\frac{\kappa}{m}} \cdot t$ ein Maß für die Frequenz der Schwingung.

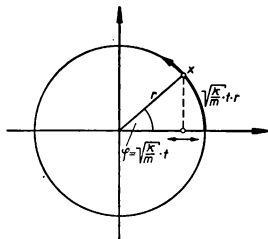


Abb. 62

Aufgaben

1. Rechnen Sie die unter a) durchgeführten Beispiele nach!
2. Welche Stromstärke fließt in dem Stromkreis der Abbildung 63, der an das Wechselstromnetz ($f = 50 \text{ Hz}$, $U = 220 \text{ V}$) angeschlossen ist?

S_1 : Selbstinduktion 2 Henry,
Ohmscher Widerstand 50 Ω

S_2 : Selbstinduktion 5 Henry,
Ohmscher Widerstand 250 Ω

S_3 : Selbstinduktion 3 Henry,
Ohmscher Widerstand 100 Ω

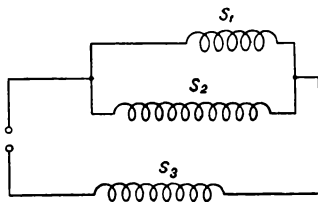


Abb. 63

15. Zusammenfassung

Wir haben die Entwicklung des Zahlenbegriffes vom Bereich der natürlichen bis zum Bereich der komplexen Zahlen verfolgt. Rückschauend können wir feststellen, daß die vorgenommenen Erweiterungen von Stufe zu Stufe durch den jeweils vorliegenden Zahlenbereich und die Bedürfnisse der Rechenpraxis vorgeschrieben sind. Wir können auch sagen, daß der Bereich der natürlichen Zahlen bereits den Ansatz zu allen Erweiterungen in sich trägt.

Bei jeder Erweiterung wurde unter Beibehaltung aller bereits gültigen Rechengesetze der Umfang der vornehmbaren Rechenoperationen erweitert, während man sich vom ursprünglichen Zahlbegriff, nämlich der Möglichkeit, mit den Zahlen zu zählen, immer mehr entfernte. Wir wollen dies noch einmal in Form einer Tabelle zusammenstellen (vgl. S. 152).

Die Tatsache, daß auch das Logarithmieren und damit alle sieben Grundrechenoperationen innerhalb der komplexen Zahlen immer durchführbar sind, können wir hier nur erwähnen. Auch hat jede konvergierende Folge komplexer Zahlen stets eine komplexe Zahl als Grenzwert, so daß auch die Probleme der Infinitesimalrechnung nicht über den Bereich der komplexen Zahlen hinausführen.

Wir haben also mit dem Bereich der komplexen Zahlen einen echten Abschluß in der Entwicklung des Zahlbegriffes erreicht. Da es keine in ihm unlösbaren Rechenaufgaben gibt, die man als „neue Zahlen“ einführen könnte, liegt auch kein praktisches Bedürfnis vor, ihn zu erweitern.

Freilich kann man auch über den Bereich der komplexen Zahlen hinausgehende abstrakte Rechenbereiche schaffen, wie Vektoren, Matrizen, Quaternionen usw. Sie können ebenso wie die komplexen Zahlen auf physikalische Probleme angewendet werden und stellen teilweise unentbehrliche Hilfsmittel zur Erforschung und Beherrschung der Natur dar. Trotzdem ist es nicht mehr angebracht, bei ihnen von Zahlen zu sprechen. Man kann nämlich beweisen, daß jeder über die komplexen Zahlen hinausgehende Rechenbereich wenigstens eines der schon für die natürlichen Zahlen geltenden Grundgesetze verletzen muß. In den meisten Fällen ist es das kommutative Gesetz der Multiplikation, welches nicht mehr aufrechterhalten werden kann.

Bereich	Vervollkommnung der Rechentechnik	Entfernung vom ursprünglichen Zahlbegriff
Bereich der natürlichen Zahlen	Addition und Multiplikation uneingeschränkt durchführbar	Abzählen im ursprünglichen Sinne des Wortes
Bereich der ganzen Zahlen	Subtraktion uneingeschränkt durchführbar	Abzählen nach zwei Seiten, etwa Thermometerskalen
Bereich der rationalen Zahlen	Division uneingeschränkt durchführbar	Statt diskreter Punktmenge dichte Punktmenge. Ermöglicht Messungen mit beliebig genauer Annäherung
Bereich der reellen Zahlen	Jede konvergierende Zahlenfolge hat einen Grenzwert. Damit Potenzieren, Radizieren und Logarithmieren weitgehend durchführbar	Kontinuierliche Punktmenge. Jeder Länge entspricht exakt ein Punkt
Bereich der komplexen Zahlen	Potenzieren, Radizieren und Logarithmieren uneingeschränkt durchführbar. Jede algebraische Gleichung n -ten Grades hat genau n Wurzeln	Lineare Anordnung muß aufgegeben werden. Dafür zweidimensionale Anordnung und Möglichkeit, Vorgänge in einer Ebene zu beschreiben

Bemerkungen: In der mittleren Spalte ist jeweils die Vervollkommnung der Rechentechnik gegenüber dem vorangehenden Bereich angegeben. Das Darüberstehende bleibt erhalten, das Darunterstehende ist noch nicht erreicht. Die Division durch Null ist stets ausgeschlossen.

IV. Zusammenstellung der wichtigsten Formeln zum Rechnen mit komplexen Zahlen

Eine komplexe Zahl kann in arithmetischer und in geometrischer Form geschrieben werden:

$$z = a + bi = r (\cos \varphi + i \cdot \sin \varphi).$$

Dabei gelten für den Realteil a und den Imaginärteil b einerseits und für den Betrag r und das Argument φ andererseits die Umrechnungsformeln:

$$a = r \cos \varphi, \quad b = r \cdot \sin \varphi,$$

$$r = +\sqrt{a^2 + b^2}, \quad \varphi = \operatorname{arctg} \frac{b}{a}.$$

Addition und Subtraktion zweier komplexer Zahlen in arithmetischer Form:

$$(a + bi) + (c + di) = (a + c) + (b + d)i$$

$$(a + bi) - (c + di) = (a - c) + (b - d)i$$

Multiplikation und Division zweier komplexer Zahlen in arithmetischer Form:

$$(a + bi) \cdot (c + di) = (ac - bd) + (ad + bc)i$$

$$\frac{a + bi}{c + di} = \frac{ac + bd}{c^2 + d^2} + \frac{bc - ad}{c^2 + d^2}i$$

Multiplikation und Division zweier komplexer Zahlen in trigonometrischer Form:

$$\begin{aligned} \{r_1 (\cos \varphi_1 + i \cdot \sin \varphi_1)\} \cdot \{r_2 (\cos \varphi_2 + i \cdot \sin \varphi_2)\} \\ = r_1 \cdot r_2 (\cos (\varphi_1 + \varphi_2) + i \cdot \sin (\varphi_1 + \varphi_2)) \end{aligned}$$

$$\begin{aligned} \{r_1 (\cos \varphi_1 + i \cdot \sin \varphi_1)\} : \{r_2 (\cos \varphi_2 + i \cdot \sin \varphi_2)\} \\ = \frac{r_1}{r_2} (\cos (\varphi_1 - \varphi_2) + i \cdot \sin (\varphi_1 - \varphi_2)) \end{aligned}$$

Für die Potenz einer komplexen Zahl in trigonometrischer Form gilt (Satz von Moivre):

$$\{r (\cos \varphi + i \cdot \sin \varphi)\}^n = r^n (\cos n\varphi + i \cdot \sin n\varphi)$$

für alle reellen (sogar komplexen) Exponenten n .

Insbesondere gilt für zwei zueinander konjugiert komplexe Zahlen:

$$(a + bi) + (a - bi) = 2a,$$

$$(a + bi) - (a - bi) = 2bi,$$

$$(a + bi)(a - bi) = a^2 + b^2.$$

INHALTSVERZEICHNIS

A. Fehlerrechnung und Näherungslösungen		Seite
I. Fehlerrechnung		3
1. Grundlegende Begriffe		3
2. Mittelwert und durchschnittlicher Fehler .		4
3. Der relative Fehler		6
4. Der Fehler eines Rechenergebnisses		8
II. Näherungslösungen		18
5. Praktische Berechnung erforderlicher Funktionswerte		18
6. Das Sekantennäherungsverfahren . .		19
7. Das Tangentennäherungsverfahren		21
B. Potenzreihen		
I. Darstellung einer Funktion durch Potenzreihen		25
1. Die Mittelwertsätze der Differentialrechnung .		25
2. Die Taylorsche Sätze		30
II. Grundbegriffe der Reihenlehre		42
3. Zahlenfolgen		42
4. Der Grenzwert einer Zahlenfolge		43
5. Reihen		45
III. Konvergenzuntersuchungen .		49
6. Einfachste Sätze über konvergente Reihen .		49
7. Konvergenzkriterien		51
8. Absolute und bedingte Konvergenz		54
IV. Potenzreihen. . . .		58
9. Allgemeines über Potenzreihen .		58
10. Die Taylorsche Reihe		60
11. Die Exponentialreihe		62
12. Die Reihen für $\sin x$ und $\cos x$		63
13. Die binomische Reihe		64
14. Die logarithmische Reihe .		68
15. Die Reihe für $\arctg x$. .		70
V. Zusammenstellung der wichtigsten Sätze und Reihenentwicklungen		76

C. Sphärische Trigonometrie

Seite

I. Grundbegriffe der Kugelgeometrie .	77
1. Großkreise und Kleinkreise, kürzeste Entfernung zweier Punkte	77
2. Das sphärische Zweieck, Winkel am Zweieck, Flächeninhalt	79
3. Das sphärische Dreieck, Flächeninhalt, Winkelsumme .	80
4. Das sphärische Dreieck und die körperliche Ecke .	82
5. Die Polarecke und das Polardreieck, die Seitensumme .	84
II. Berechnung des sphärischen Dreiecks	86
6. Das rechtwinklige sphärische Dreieck, die Nepersche Regel	86
7. Das schiefwinklige Dreieck. .	88
8. Die sechs Fälle der Dreiecksberechnung	91
III. Anwendung auf die Erdkugel .	95
9. Das Gradnetz und die Größe der Erde	95
10. Das Poldreieck, die Orthodrome und die Kurswinkel	96
11. Loxodrome, Scheitelpunkt, Schnitt mit Meridian- und Parallelkreisen .	98

D. Komplexe Zahlen

I. Vom Bereich der natürlichen Zahlen zum Bereich der reellen Zahlen .	109
1. Der Bereich der natürlichen Zahlen	110
2. Die Erweiterung von Zahlenbereichen	114
3. Der Bereich der rationalen Zahlen	117
4. Der Bereich der reellen Zahlen	119
II. Der Bereich der komplexen Zahlen .	121
5. Einführung der imaginären Zahlen	122
6. Das Rechnen mit imaginären Zahlen	124
7. Einführung der komplexen Zahlen .	126
8. Die Rechenoperationen erster und zweiter Stufe im Bereich der komplexen Zahlen	127
9. Die Gaußsche Zahlenebene .	131
III. Rechenoperationen im Bereich der komplexen Zahlen	134
10. Die trigonometrische Form der komplexen Zahlen	134
11. Multiplikation und Division komplexer Zahlen in trigonometrischer Form	135
12. Potenzen und Wurzeln komplexer Zahlen, der Satz von Moivre	138
13. Algebraische Gleichungen	141
14. Anwendung komplexer Zahlen in Physik und Technik	147
15. Zusammenfassung	151
IV. Zusammenstellung der wichtigsten Formeln zum Rechnen mit komplexen Zahlen	153

