

Analysis I

Vorlesung 1

Mengen



Georg Cantor (1845-1918) ist der Schöpfer der Mengentheorie.



David Hilbert (1862-1943) nannte sie ein *Paradies*, aus dem die Mathematiker nie mehr vertrieben werden dürften.

Eine *Menge* ist eine Ansammlung von wohlunterschiedenen Objekten, die die *Elemente* der Menge heißen. Mit „wohlunterschieden“ meint man, dass es klar ist, welche Objekte als gleich und welche als verschieden angesehen werden. Die *Zugehörigkeit* eines Elementes x zu einer Menge M wird durch

$$x \in M$$

ausgedrückt, die Nichtzugehörigkeit durch

$$x \notin M.$$

Für jedes Element(symbol) gilt stets genau eine dieser zwei Möglichkeiten.

Für Mengen gilt das *Extensionalitätsprinzip*, d.h. eine Menge ist durch die in ihr enthaltenen Elemente eindeutig bestimmt, darüber hinaus bietet sie keine Information. Insbesondere stimmen zwei Mengen überein, wenn beide die gleichen Elemente enthalten.

Die Menge, die kein Element besitzt, heißt *leere Menge* und wird mit

$$\emptyset$$

bezeichnet.

Eine Menge N heißt *Teilmenge* einer Menge M , wenn jedes Element aus N auch zu M gehört. Man schreibt dafür

$$N \subseteq M$$

(manche schreiben dafür $N \subset M$). Man sagt dafür auch, dass eine *Inklusion* $N \subseteq M$ vorliegt. Im Nachweis, dass $N \subseteq M$ ist, muss man zeigen, dass für ein beliebiges Element $x \in N$ ebenfalls die Beziehung $x \in M$ gilt.¹ Dabei darf man lediglich die Eigenschaft $x \in N$ verwenden.

Aufgrund des Extensionalitätsprinzips hat man das folgende wichtige *Gleichheitsprinzip für Mengen*, dass

$$M = N \text{ genau dann, wenn } N \subseteq M \text{ und } M \subseteq N$$

gilt. In der mathematischen Praxis bedeutet dies, dass man die Gleichheit von zwei Mengen dadurch nachweist, dass man (in zwei voneinander unabhängigen Teilargumentationen) die beiden Inklusionen zeigt. Dies hat auch den kognitiven Vorteil, dass das Denken eine Zielrichtung bekommt, dass klar die Voraussetzung, die man verwenden darf, von der gewünschten Schlussfolgerung, die man aufzeigen muss, getrennt wird. Hier spiegelt sich das aussagenlogische Prinzip wieder, dass die Äquivalenz von zwei Aussagen die wechselseitige Implikation bedeutet, und durch den Beweis der beiden einzelnen Implikationen bewiesen wird.

Beschreibungsmöglichkeiten für Mengen

Es gibt mehrere Möglichkeiten, eine Menge anzugeben. Die einfachste ist, die zu der Menge gehörenden Elemente aufzulisten, wobei es auf die Reihenfolge der Elemente nicht ankommt. Bei endlichen Mengen ist dies unproblematisch, bei unendlichen Mengen muss man ein „Bildungsgesetz“ für die Elemente angeben.

Die wichtigste Menge, die man zumeist als eine fortgesetzte Auflistung einführt, ist die Menge der natürlichen Zahlen

$$\mathbb{N} = \{0, 1, 2, 3, \dots\}.$$

Hier wird eine bestimmte Zahlenmenge durch die Anfangsglieder von erlaubten Zifferfolgen angedeutet. Wichtig ist, dass mit $\mathbb{N}$ nicht eine Menge von bestimmten Ziffern gemeint ist, sondern die durch die Ziffern repräsentierten

¹In der Sprache der Quantorenlogik kann man eine Inklusion verstehen als die Aussage $\forall x(x \in N \rightarrow x \in M)$.

Zahlwerte. Eine natürliche Zahl hat viele Darstellungsarten, die Ziffernrepräsentation im Zehnersystem ist nur eine davon, wenn auch eine besonders übersichtliche.

Wir besprechen Mengenbeschreibung durch Eigenschaften. Es sei eine Menge M gegeben. In ihr gibt es gewisse Elemente, die gewisse Eigenschaften E (Prädikate) erfüllen können oder aber nicht. Zu einer Eigenschaft E gehört innerhalb von M die Teilmenge bestehend aus allen Elementen aus M , die diese Eigenschaft erfüllen. Man beschreibt eine durch eine Eigenschaft definierte Teilmenge meist als

$$\{x \in M \mid E(x)\} = \{x \in M \mid x \text{ besitzt die Eigenschaft } E\}.$$

Dies geht natürlich nur mit solchen Eigenschaften, für die die Aussage $E(x)$ eine wohldefinierte Bedeutung hat. Dieser Konstruktion entspricht in der Alltagssprache eine Formulierung mit einem Relativsatz, im Sinne von diejenigen Objekte, auf die die Eigenschaft E zutrifft. Wenn man eine solche Teilmenge einführt, so gibt man ihr häufig sofort einen Namen (in dem auf die Eigenschaft E Bezug genommen werden kann, aber nicht muss). Z.B. kann man einführen

$$G = \{x \in \mathbb{N} \mid x \text{ ist gerade}\},$$

$$U = \{x \in \mathbb{N} \mid x \text{ ist ungerade}\},$$

$$Q = \{x \in \mathbb{N} \mid x \text{ ist eine Quadratzahl}\}$$

$$\mathbb{P} = \{x \in \mathbb{N} \mid x \text{ ist eine Primzahl}\}.$$

Für die Mengen in der Mathematik sind meist eine Vielzahl an mathematischen Eigenschaften relevant und daher gibt es meist auch eine Vielzahl an relevanten Teilmengen. Aber auch bei alltäglichen Mengen, wie etwa die Menge K der Studierenden in einem Kurs, gibt es viele wichtige Eigenschaften, die gewisse Teilmengen festlegen, wie etwa

$$O = \{x \in K \mid x \text{ kommt aus Osnabrück}\},$$

$$P = \{x \in K \mid x \text{ studiert im Nebenfach Physik}\},$$

$$D = \{x \in K \mid x \text{ hat im Dezember Geburtstag}\}.$$

Die Menge K ist dabei selbst durch eine Eigenschaft festgelegt, es ist ja

$$K = \{x \mid x \text{ ist Studierender in diesem Kurs}\}.$$

Mengenoperationen

So, wie man Aussagen zu neuen Aussagen verknüpfen kann, gibt es Operationen, mit denen aus Mengen neue Mengen entstehen.

- *Vereinigung*

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

- *Durchschnitt*

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

- *Differenzmenge*

$$A \setminus B := \{x \mid x \in A \text{ und } x \notin B\}.$$

Diese Operationen ergeben nur dann einen Sinn, wenn die beteiligten Mengen als Teilmengen in einer gemeinsamen Grundmenge gegeben sind. Dies sichert, dass man über die gleichen Elemente spricht. Häufig wird diese Grundmenge nicht explizit angegeben, dann muss man sie aus dem Kontext erschließen. Ein Spezialfall der Differenzmenge bei einer gegebenen Grundmenge ist das *Komplement* einer Teilmenge $A \subseteq G$, das durch

$$\complement A := G \setminus A = \{x \in G \mid x \notin A\}$$

definiert ist. Wenn zwei Mengen einen leeren Schnitt haben, also $A \cap B = \emptyset$ gilt, so nennen wir sie *disjunkt*.

Konstruktion von Mengen

Die meisten Mengen in der Mathematik ergeben sich ausgehend von einigen wenigen Mengen wie beispielsweise den endlichen Mengen und $\mathbb{N}$ durch bestimmte Konstruktionen von neuen Mengen aus schon bekannten oder schon zuvor konstruierten Mengen.² Wir definieren:³

²Darunter fallen auch der Schnitt und die Vereinigung, doch bleiben diese innerhalb einer vorgegebenen Grundmenge, während es hier um Konstruktionen geht, die darüber hinaus gehen.

³Definitionen werden in der Mathematik zumeist als solche deutlich herausgestellt und bekommen eine Nummer, damit man auf sie einfach Bezug nehmen kann. Es wird eine Situation beschrieben, bei der die verwendeten Begriffe schon zuvor definiert worden sein mussten, und in dieser Situation wird einem neuen Konzept ein Name (eine Bezeichnung) gegeben. Dieser Name wird *kursiv* gesetzt. Man beachte, dass das Konzept auch ohne den neuen Namen formulierbar ist, der neue Name ist nur eine Abkürzung für das Konzept. Sehr häufig hängen die Begriffe von Eingaben ab, wie den beiden Mengen in dieser Definition. Bei der Namensgebung herrscht eine gewisse Willkür, so dass die Bedeutung der Bezeichnung im mathematischen Kontext sich allein aus der expliziten Definition, aber nicht aus der alltäglichen Wortbedeutung erschließen lässt.

DEFINITION 1.1. Es seien zwei Mengen L und M gegeben. Dann nennt man die Menge

$$L \times M = \{(x, y) \mid x \in L, y \in M\}$$

die *Produktmenge*⁴ der beiden Mengen.

Die Elemente der Produktmenge nennt man *Paare* und schreibt (x, y) . Dabei kommt es wesentlich auf die Reihenfolge an. Die Produktmenge besteht also aus allen Paarkombinationen, wo in der ersten *Komponente* ein Element der ersten Menge und in der zweiten Komponente ein Element der zweiten Menge steht. Zwei Paare sind genau dann gleich, wenn sie in beiden Komponenten gleich sind.

Bei einer Produktmenge können natürlich auch beide Mengen gleich sein, beispielsweise ist $\mathbb{R} \times \mathbb{R}$ die reelle Ebene. In diesem Fall ist es verlockend, die Reihenfolge zu verwechseln, und also besonders wichtig, darauf zu achten, dies nicht zu tun. Wenn es in der ersten Menge n Elemente und in der zweiten Menge k Elemente gibt, so gibt es in der Produktmenge $n \cdot k$ Elemente. Wenn eine der beiden Mengen leer ist, so ist auch die Produktmenge leer. Man kann auch für mehr als nur zwei Mengen die Produktmenge bilden, worauf wir bald zurückkommen werden.

BEISPIEL 1.2. Es sei V die Menge aller Vornamen (sagen wir der Vornamen, die in einer bestimmten Grundmenge an Personen wirklich vorkommen) und N die Menge aller Nachnamen. Dann ist

$$V \times N$$

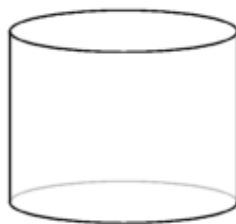
die Menge aller Namen. Elemente davon sind in Paarschreibweise beispielsweise (Heinz, Müller), (Petra, Müller) und (Lucy, Sonnenschein). Aus einem Namen lässt sich einfach der Vorname und der Nachname herauslesen, indem man entweder auf die erste oder auf die zweite Komponente des Namens schaut. Auch wenn alle Vornamen und Nachnamen für sich genommen vorkommen, so muss natürlich nicht jeder daraus gebastelte mögliche Name wirklich vorkommen. Bei der Produktmenge werden eben alle Kombinationsmöglichkeiten aus den beiden beteiligten Mengen genommen.

BEISPIEL 1.3. Ein Schachbrett (genauer: die Menge der Felder auf einem Schachbrett, auf denen eine Figur stehen kann) ist die Produktmenge $\{a, b, c, d, e, f, g, h\} \times \{1, 2, 3, 4, 5, 6, 7, 8\}$. Jedes Feld ist ein Paar, beispielsweise $(a, 1)$, $(d, 4)$, $(c, 7)$. Da die beteiligten Mengen verschieden sind, kann man statt der Paarschreibweise einfach $a1$, $d4$, $c7$ schreiben. Diese Notation ist der Ausgangspunkt für die Beschreibung von Stellungen und von ganzen Partien.

⁴Man spricht auch vom *kartesischen Produkt* der beiden Mengen.



Wenn zwei geometrische Punktmengen A und B gegeben sind, beispielsweise als Teilmengen einer Ebene E , so kann man die Produktmenge $A \times B$ als Teilmenge von $E \times E$ auffassen. Dadurch entsteht ein neues geometrisches Gebilde, das man manchmal auch in einer kleineren Dimension realisieren kann.



Ein Zylindermantel ist die Produktmenge aus einem Kreis und einer Strecke.

BEISPIEL 1.4. Es sei S ein Kreis, worunter wir die Kreislinie verstehen, und I eine Strecke. Der Kreis ist eine Teilmenge einer Ebene E und die Strecke ist eine Teilmenge einer Geraden G , so dass für die Produktmenge die Beziehung

$$S \times I \subseteq E \times G$$

gilt. Die Produktmenge $E \times G$ stellt man sich als einen dreidimensionalen Raum vor, und darin ist die Produktmenge $S \times I$ ein Zylindermantel.

Eine andere wichtige Konstruktion, um aus einer Menge eine neue Menge zu erhalten, ist die Potenzmenge.

DEFINITION 1.5. Zu einer Menge M nennt man die Menge aller Teilmengen von M die *Potenzmenge* von M . Sie wird mit

$$\mathfrak{P}(M)$$

bezeichnet.

Es ist also

$$\mathfrak{P}(M) = \{T \mid T \text{ ist Teilmenge von } M\}.$$

Wenn M die Menge der Kursteilnehmer ist, so kann man sich jede Teilmenge als eine kursinterne Party vorstellen, zu der eine gewisse Auswahl an Leuten hingeht (es werden also die Parties mit den anwesenden Leuten identifiziert). Die Potenzmenge ist dann die Menge aller möglichen Parties. Wenn eine Menge n Elemente besitzt, so besitzt ihre Potenzmenge 2^n Elemente.

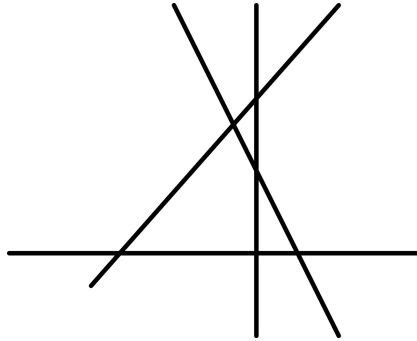
Induktion

Die natürlichen Zahlen sind dadurch ausgezeichnet, dass man jede natürliche Zahl ausgehend von der 0 durch den Zählprozess (das sukzessive Nachfolgernehmen) erreichen kann. Daher können mathematische Aussagen, die von natürlichen Zahlen abhängen, mit dem Beweisprinzip der *vollständigen Induktion* bewiesen werden. Das folgende Beispiel soll an dieses Argumentationsschema herañführen.

BEISPIEL 1.6. Wir betrachten in der Ebene E eine Konfiguration von n Geraden und fragen uns, was die maximale Anzahl an Schnittpunkten ist, die eine solche Konfiguration haben kann. Dabei ist es egal, ob wir uns die Ebene als einen $\mathbb{R}^2$ (eine kartesische Ebene mit Koordinaten) oder einfach elementargeometrisch vorstellen, wichtig ist im Moment allein, dass sich zwei Geraden in genau einem Punkt schneiden können oder aber parallel sein können. Wenn n klein ist, so findet man relativ schnell die Antwort.

n	0	1	2	3	4	5	n
$S(n)$	0	0	1	3	6	?	?

Doch schon bei etwas größerem n ($n = 5, 10, \dots$?) kann man ins Grübeln kommen, da man sich die Situation irgendwann nicht mehr präzise vorstellen kann. Aus einer präzisen Vorstellung wird eine Vorstellung von vielen Geraden mit vielen Schnittpunkten, woraus man aber keine exakte Anzahl der Schnittpunkte ablesen kann. Ein sinnvoller Ansatz zum Verständnis des Problems ist es, sich zu fragen, was eigentlich passiert, wenn eine neue Gerade hinzukommt, wenn also aus n Geraden $n+1$ Geraden werden. Angenommen, man weiß aus irgendeinem Grund, was die maximale Anzahl der Schnittpunkte bei n Geraden ist, im besten Fall hat man dafür eine Formel. Wenn man dann versteht, wie viele neue Schnittpunkte maximal bei der Hinzunahme von einer neuen Geraden hinzukommen, so weiß man, wie die Anzahl der maximalen Schnittpunkte von $n+1$ Geraden lautet.



Dieser Übergang ist in der Tat einfach zu verstehen. Die neue Gerade kann höchstens jede der n alten Geraden in genau einem Punkt schneiden, deshalb kommen höchstens n neue Schnittpunkte hinzu. Wenn man die neue Gerade so wählt, dass sie zu keiner der gegebenen Geraden parallel ist (was möglich ist, da es unendlich viele Richtungen gibt) und ferner so wählt, dass die neuen Schnittpunkte von den schon gegebenen Schnittpunkten der Konfiguration verschieden sind (was man erreichen kann, indem man die neue Gerade parallel verschiebt, um den alten Schnittpunkten auszuweichen), so erhält man genau n neue Schnittpunkte. Von daher ergibt sich die (vorläufige) Formel

$$S(n+1) = 1 + 2 + 3 + \cdots + n - 2 + n - 1 + n$$

bzw.

$$S(n) = 1 + 2 + 3 + \cdots + n - 3 + n - 2 + n - 1,$$

also einfach die Summe der ersten $n - 1$ natürlichen Zahlen.

Im vorstehenden Beispiel liegt eine Summe vor, wobei die Anzahl der Summanden selbst variieren kann. Für eine solche Situation ist das *Summenzeichen* sinnvoll. Für gegebene reelle Zahlen $a_1, \dots, a_n$ bedeutet

$$\sum_{k=1}^n a_k := a_1 + a_2 + \cdots + a_{n-1} + a_n.$$

Dabei hängen im Allgemeinen die a_k in einer formelhaften Weise von k ab, beispielsweise ist im Beispiel $a_k = k$, es könnte aber auch etwas wie $a_k = 2k + 1$ oder $a_k = k^2$ vorliegen. Der k -te Summand der Summe ist jedenfalls a_k , dabei nennt man k den *Index* des Summanden. Entsprechend ist das *Produktzeichen* definiert, nämlich durch

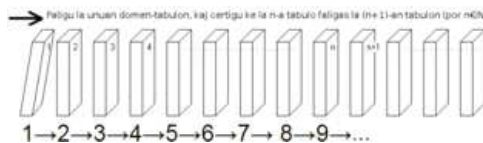
$$\prod_{k=1}^n a_k := a_1 \cdot a_2 \cdots a_{n-1} \cdot a_n.$$

BEISPIEL 1.7. Wir möchten für die Summe der ersten n Zahlen, die die maximale Anzahl der Schnittpunkte in einer Konfiguration aus $n - 1$ Geraden

angibt, eine einfachere Formel angeben. Und zwar behaupten wir, dass

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

Für kleinere Zahlen n stimmt dies aus dem einfachen Grund, dass links und rechts dasselbe herauskommt. Um die Gleichung allgemein zu beweisen, überlegen wir uns, was links und was rechts passiert, wenn wir das n um 1 erhöhen, so wie wir zuvor die Geradenkonfiguration um eine zusätzliche Gerade verkompliziert haben. Auf der linken Seite kommt einfach der zusätzliche Summand $n+1$ hinzu. Auf der rechten Seite haben wir den Übergang von $\frac{n(n+1)}{2}$ nach $\frac{(n+1)(n+1+1)}{2}$. Wenn wir zeigen können, dass die Differenz zwischen diesen beiden Brüchen ebenfalls $n+1$ ist, so verhält sich die rechte Seite genauso wie die linke Seite. Dann kann man so schließen: die Gleichung gilt für die kleinen n , etwa für $n = 1$. Durch den Differenzenvergleich gilt es auch für das nächste n , also für $n = 2$, durch den Differenzenvergleich gilt es für das nächste n , u.s.w. Da dieses Argument immer funktioniert, und da man jede natürliche Zahl irgendwann durch sukzessives Nachfolgernehmen erreicht, gilt die Formel für jede natürliche Zahl.



Eine Visualisierung des Induktionsprinzips. Wenn die Steine nah beieinander stehen und der erste umgestoßen wird, so fallen alle Steine um.

Die folgende Aussage begründet das Prinzip der vollständigen Induktion.

SATZ 1.8. *Für jede natürliche Zahl n sei eine Aussage $A(n)$ gegeben. Es gelte*

- (1) $A(0)$ ist wahr.
- (2) Für alle n gilt: wenn $A(n)$ gilt, so ist auch $A(n+1)$ wahr.

Dann gilt $A(n)$ für alle n .

Beweis. Wegen der ersten Voraussetzung gilt $A(0)$. Wegen der zweiten Voraussetzung gilt auch $A(1)$. Deshalb gilt auch $A(2)$. Deshalb gilt auch $A(3)$. Da man so beliebig weitergehen kann und dabei jede natürliche Zahl erhält, gilt die Aussage $A(n)$ für jede natürliche Zahl n . $\square$

Der Nachweis von $A(0)$ heißt dabei der *Induktionsanfang* und der Schluss von $A(n)$ auf $A(n+1)$ heißt der *Induktionsschritt*. Innerhalb des Induktionsschrittes nennt man die Gültigkeit von $A(n)$ die *Induktionsvoraussetzung*. In manchen Situationen ist die Aussage $A(n)$ erst für $n \geq n_0$ für ein gewisses n_0

(definiert oder) wahr. Dann beweist man im Induktionsanfang die Aussage $A(n_0)$ und den Induktionsschluss führt man für $n \geq n_0$ durch.

Wir begründen nun die Gleichheit

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

mit dem Induktionsprinzip.

Beim Induktionsanfang ist $n = 1$, daher besteht die Summe links nur aus einem Summanden, nämlich der 1, und daher ist die Summe 1. Die rechte Seite ist $\frac{1 \cdot 2}{2} = 1$, so dass die Formel für $n = 1$ stimmt.

Für den Induktionsschritt setzen wir voraus, dass die Formel für ein $n \geq 1$ gilt, und müssen zeigen, dass sie dann auch für $n+1$ gilt. Dabei ist n beliebig. Es ist

$$\begin{aligned} \sum_{k=1}^{n+1} k &= \left(\sum_{k=1}^n k \right) + n + 1 \\ &= \frac{n(n+1)}{2} + n + 1 \\ &= \frac{n(n+1) + 2(n+1)}{2} \\ &= \frac{(n+2)(n+1)}{2}. \end{aligned}$$

Dabei haben wir für die zweite Gleichheit die Induktionsvoraussetzung verwendet. Der zuletzt erhaltene Term ist die rechte Seite der Formel für $n+1$, also ist die Formel bewiesen.

BEMERKUNG 1.9. Induktionsbeweise kommen ständig vor. Die Voraussetzung dafür, dass man diese Beweismethode anwenden kann, ist, dass ein Aussagenschema vorliegt, das von einer (variablen) natürlichen Zahl n abhängt. Diese natürliche Zahl n nennt man auch die *Induktionsvariable*, über die Induktion geführt wird. In der Aussage können ansonsten beliebige andere mathematische Objekte vorkommen und die natürliche Zahl n kann jeweils sehr unterschiedliche Bedeutungen haben. Sie kann für den Exponenten einer reellen Zahl (siehe beispielsweise Satz 3.9 oder Satz 4.10), den Grad eines Polynoms (etwa Satz 11.3), den Differenzierbarkeitsgrad (etwa Aufgabe 19.15) oder die Anzahl von Vektoren (etwa Lemma 22.3 (Lineare Algebra (Osna-brück 2017-2018))) stehen.

Abbildungsverzeichnis

Quelle = Georg Cantor 1894.jpg , Autor = Benutzer Taxiarchos228 auf Commons, Lizenz = PD	1
Quelle = David Hilbert 1886.jpg , Autor = Unbekannt (1886), Lizenz = PD	1
Quelle = Chess board blank.svg , Autor = Benutzer Beao auf Commons, Lizenz = CC-by-sa 3.0	6
Quelle = Geometri cylinder.png , Autor = Benutzer Anp auf sv Wikipedia, Lizenz = PD	6
Quelle = 4Geraden6Schnittpunkte.png , Autor = Benutzer Mgausmann auf CC-by-sa 4.0, Lizenz =	8
Quelle = Domen-indukto.gif , Autor = Joachim Mohr, Lizenz = CC-by-sa 3.0	9
Erläuterung: Die in diesem Text verwendeten Bilder stammen aus Commons (also von http://commons.wikimedia.org) und haben eine Lizenz, die die Verwendung hier erlaubt. Die Bilder werden mit ihren Dateinamen auf Commons angeführt zusammen mit ihrem Autor bzw. Hochlader und der Lizenz.	11
Lizenzklärung: Diese Seite wurde von Holger Brenner alias Bocardodarapti auf der deutschsprachigen Wikiversity erstellt und unter die Lizenz CC-by-sa 3.0 gestellt.	11