

Molekularbiologie der Zelle

Michael Kubi

Inhalt

Kapitel 1: Einführung	7
Was ist Leben?	7
Der Aufbau von Zellen – ein Überblick	9
Kapitel 2: Atome, Periodensystem und Atombindung	13
Atome – Grundaufbau	13
Periodensystem der Elemente (PSE)	15
Ionen und Ionenbindung.....	17
Atombindung	18
Isotope.....	20
Wasser und pH-Wert.....	20
Kapitel 3: Proteine, Zucker und Fette.....	26
Wir sind organisch!	26
Aminosäuren.....	30
Proteine.....	34
Kohlenhydrate	37
Lipide.....	40
Kapitel 4: chemische Reaktion, Katalyse, Enzyme und ATP.....	43
Chemische Reaktion, Oxidation & Reduktion	43
Exotherm, endotherm und Aktivierungsenergie.....	44
Enzyme.....	46
Enzymkinetik.....	48
ATP	50
Kapitel 5: Biomembranen.....	52
Flüssig-Mosaik-Modell	52
Semipermeable Membran und Diffusion	53
Erleichterte Diffusion und Aktiver Stofftransport.....	56
Kapitel 6: Zellatmung	58
Glykolyse	59
Gärung und Gluconeogenese.....	62
Oxidative Decarboxylierung von Pyruvat.....	63
Citratzyklus	64
Atmungskette.....	67
Endbilanz	71
Kapitel 7: Photosynthese	73

Chlorophyll und Chloroplasten	74
Lichtreaktion	78
Bildung von ATP	82
Zyklischer Elektronentransport	83
Calvin Zyklus.....	83
CO ₂ – Fixierungsphase	84
Reduktionsphase	85
Regenerationsphase	85
Calvin Zyklus Bilanz	86
CO₂, Photosynthese, Pflanzenwachstum und Klimawandel	86
Photorespiration.....	88
Photosynthese Typen	90
C3 Pflanzen	90
C4 Pflanzen	90
CAM Pflanzen	93
Weitere Photosynthese-Faktoren	95
Abhängigkeit von der CO ₂ -Konzentration	95
Chlorophyllgehalt	96
Licht.....	96
Temperatur	98
Wasser und Mineralstoffe	98
Literatur zur Ökologie der Photosynthese	99
Kapitel 8: DNA, Chromosomen, Genom	100
Historisches	100
Nukleinsäuren.....	103
RNA	105
Organisation der Genome	105
Das menschliche Genom.....	109
Kapitel 9: DNA-Replikation.....	112
Meselson-Stahl-Experiment.....	112
Ablauf der Replikation.....	114
Telomere und Alterungsprozess	115
DNA-Reparatur.....	116
Kapitel 10: Von der DNA zum Protein	119
Ein-Gen-Ein-Enzym-(Polypeptid)-Hypothese.....	119

Genexpression und zentrales Dogma der Molekularbiologie	121
Transkription.....	122
Initiation.....	122
Elongation:	123
Termination:.....	124
Genetischer Code	124
Prozessierung & Splicing	126
Translation.....	128
Initiation:.....	131
Elongation:	131
Termination:.....	132
Kapitel 11: Genregulation und Epigenetik	133
Genregulation bei Prokaryoten	134
Genregulation bei Eukaryoten.....	137
Epigenetik.....	140
Alternatives Spleißen und RNA-Interferenz	144
Kapitel 12: Zellkommunikation.....	149
Liganden	150
Rezeptoren	152
Signalkaskaden und sekundäre Botenstoffe.....	157
Kapitel 13: Mitose und Zellzyklus	162
Zellzyklus.....	162
Mitose	166
Spindelapparat.....	168
Cytokinese.....	170
Kapitel 14: Meiose, Rekombination und Kernphasenwechsel.....	171
Meiose	171
Typen von Keimzellen	174
Crossing Over.....	176
Kernphasenwechsel	177
Haplonten.....	178
Diplonten.....	178
Diplo-Haplonten	179
Kapitel 15: Endoplasmatisches Reticulum, Golgi Apparat, Lysosom	180
Endoplasmatisches Reticulum (ER) und Golgi-Apparat (GA).....	180

Proteinmodifikation und -transport	181
Funktionen des Golgi-Apparates	184
Rolle des Cytoskeletts beim Transport.....	186
Lysosomen.....	188
Peroxisomen	190
Kapitel 16: Zellgewebe und Zellverknüpfung	192
Gewebetypen bei Tierzellen	192
Oberflächenstrukturen und extrazelluläre Matrix	193
Gap Junctions.....	195
Desmosomen	197
Tight Junctions.....	199
Mikrovilli und Mikrofilamente.....	200
Mikrofilamente und Muskelgewebe	204
Pflanzengewebe	206
Plasmodesmen	208
Kapitel 17: Zelltod und Krebs	210
Krebs	211

Kapitel 1: Einführung

Mit diesem Artikel starte ich eine Artikelserie, die sich mit dem Aufbau und Funktionen der Zelle befasst. Die Zellbiologie, auch Cytologie genannt, ist ein Teilgebiet der Biologie, welches die biologischen Vorgänge auf zellulärer Ebene zu verstehen versucht. Die Zellbiologie hat engste Kontakte mit anderen biologischen Disziplinen, wie der Biochemie, Molekularbiologie, Zoologie und Botanik. Mein Wunsch ist es verschiedenste Prozesse, die in einer Zelle geschehen, darzustellen. Ihr erfahrt etwas über den Aufbau der Zelle, über Zellatmung, Photosynthese und andere Stoffwechselprozesse sowie die Vorgänge der Molekulargenetik, wie Aufbau und Funktion der DNA, die Proteinbiosynthese, Genregulation etc.

Was ist Leben?

Unser Planet hat eine schier unendliche Anzahl verschiedenster Lebensformen, die die Materie aus der Umgebung aufnehmen und diese nicht nur dazu nutzen, um sich selbst zu erhalten, sondern auch Kopien von sich anfertigen, sich also vermehren. Doch Leben als solches zu definieren ist tatsächlich nicht einfach. Denn wenn es um das Leben geht, denken wir häufig an uns selbst oder vielleicht an unsere Haustiere. Doch auch Pflanzen leben, Pilze und Bakterien auch. Viren gelten nicht als Lebewesen, stehen aber in ihrer Daseinsform zwischen belebter und unbelebter Materie.

Was aber vereint alle Lebewesen? Wir können folgende Eigenschaften charakterisieren.

1. Lebewesen haben einen Stoffwechsel, auch Metabolismus genannt. Der Stoffwechsel ist die Gesamtheit der chemischen Prozesse in einem Lebewesen, die als Folge zur Umwandlung von Stoffen führt und steht damit für die Aufnahme, den Transport und die chemische Umwandlung von Stoffen in einem Organismus sowie die Abgabe von Stoffwechselendprodukten an die Umgebung. Diese biochemischen Vorgänge dienen dem Aufbau und der Erhaltung der Körpersubstanz (Baustoffwechsel) sowie der Energiegewinnung (Energistoffwechsel) und damit der Aufrechterhaltung der Körperfunktionen.
2. Reizbarkeit ist ebenfalls ein Kennzeichen von Lebewesen. Lebewesen können über Rezeptoren physikalische und chemische Reize aus der Umwelt empfangen und entsprechend ihrer artspezifischen Zusammensetzung und Struktur darauf reagieren.
3. Lebewesen sind in der Lage sich selbst aufrechtzuerhalten. Man spricht von Homöostase. Homöostase erzeugt ein dynamisches Gleichgewicht und ist damit ein essenzielles Prinzip für die Lebenserhaltung und Funktion eines Organismus. Beispiele hierfür sind u. a. Regelung des Kreislaufs, der Körpertemperatur, des pH-Wertes, des Wasser- und Elektrolythaushaltes oder die Steuerung des Hormonhaushaltes.
4. Fortpflanzung und Vererbung sind ein weiteres Merkmal von Lebewesen. Die Fortpflanzung stellt sicher, dass Individuen einer neuen Generation entstehen. Man unterscheidet dabei zwischen asexueller und sexueller Vermehrung. Die Fortpflanzungsfähigkeit der meisten Lebewesen beruht auf der Fähigkeit der

Zelle zu teilen (man spricht von Mitose). Bei einzelligen Lebewesen bedeutet jede Zellteilung eine (asexuelle) Fortpflanzung. Sexuelle Fortpflanzung beruht auf der Verschmelzung von Geschlechtszellen (Gameten) und erlaubt Neukombination von Erbanlagen.

5. Lebende Systeme sind durch das funktionale Zusammenwirken von Nucleinsäuren und Proteinen gekennzeichnet. Alle Lebewesen benutzen doppelsträngige Desoxyribonucleinsäuren (DNA) als genetisches Material. Die Nucleinsäuren dienen dabei als Vorlage für ihre eigene Synthese und für die Synthese der verschiedenen Proteine. Diese sorgen u.a. durch Stoffwechsel und den Aufbau der Lebewesen.
6. Die Nachkommen einer Generation sind, vor allem bei der sexuellen Vermehrung, nicht mit ihren Eltern identisch. Es gibt also Variation, hervorgerufen durch Mutationen der DNA oder durch Rekombination der Gene durch sexuelle Fortpflanzung und anderen Mechanismen. Variationen sind die Grundlage der Evolution, denn durch verschiedene Umweltbedingungen sind bestimmte Variationen vorteilhafter. Individuen mit diesen vorteilhaften Variationen überleben und vererben ihre Eigenschaften an ihre Nachkommen. Leben erzeugt also nicht nur sich selbst, sondern verändert sich. Evolution ist somit ein ständiger Begleiter des Lebens.
7. Und das für diese Videoreihe wichtigste Merkmal des Lebens ist: alle Lebewesen bestehen aus Zellen. Die Zelle ist bei allen Lebewesen die kleinste selbst reproduzierende Einheit.

Diese Eigenschaften von Leben haben keinen Anspruch auf Vollständigkeit, sollen aber für diese Videoreihe genügen. Fallen euch noch andere Eigenschaften ein? Dann schreibt sie in die Kommentare.

Der Aufbau von Zellen – ein Überblick

Wir betrachten uns die Organisation der Zelle nun etwas genauer. Wir kennen 2 Grundformen der Zellorganisation: die Procyte und die Eucyte (Abb. 1)

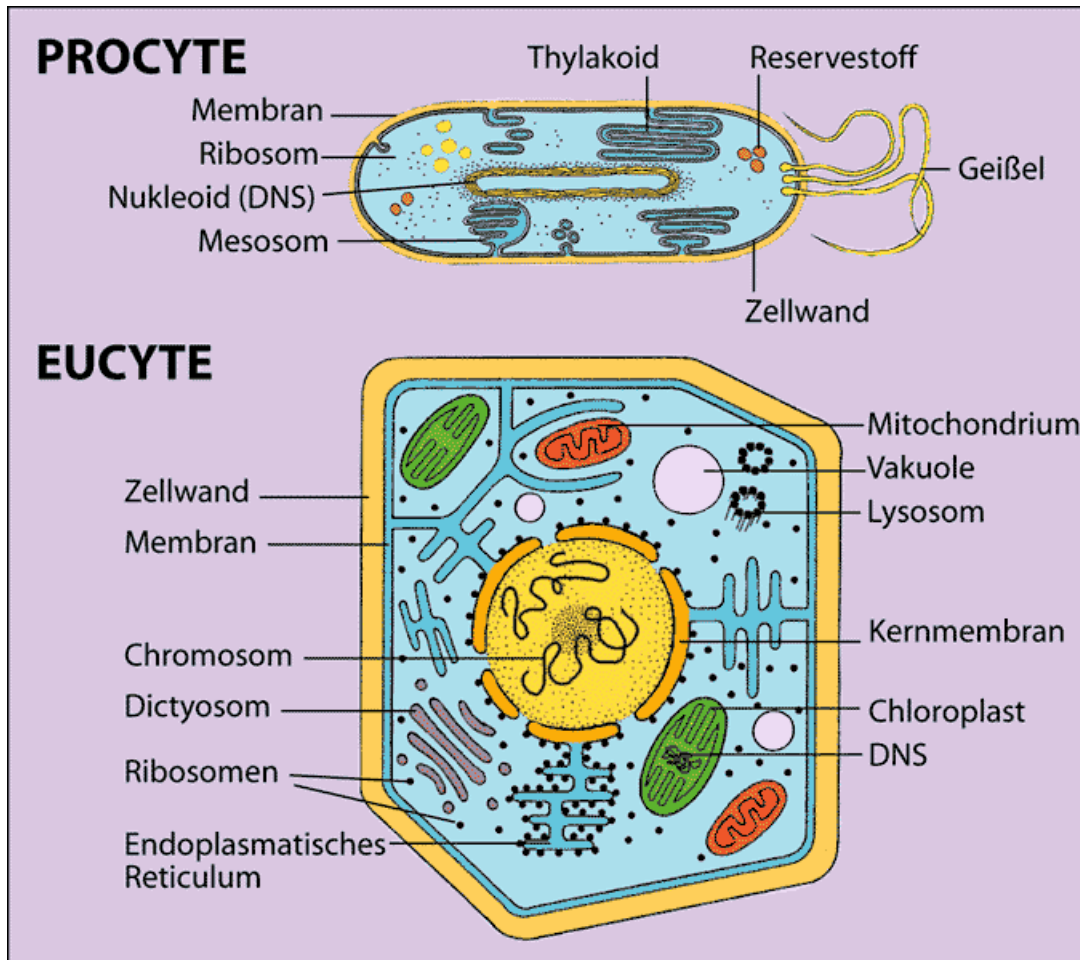


Abbildung 1 Procyte und Eucyte

Die Procyte findet sich bei den Prokaryoten und stellt die ursprünglichste Organismenform dar. Zu den Prokaryoten gehören die Bakterien und Archaeobakterien.

Die Eucyte findet sich bei den sog. Eukaryoten und findet sich bei allen Tieren, Pflanzen und Pilzen.

Der auffälligste Unterschied zwischen der Procyte und Eucyte liegt im Fehlen bzw. Vorhandensein eines Zellkerns. Procyten haben keinen Zellkern, ihre DNA liegt frei im Zellplasma. Eucyten hingegen haben einen Zellkern, in der sich ihre Erbinformation befindet.

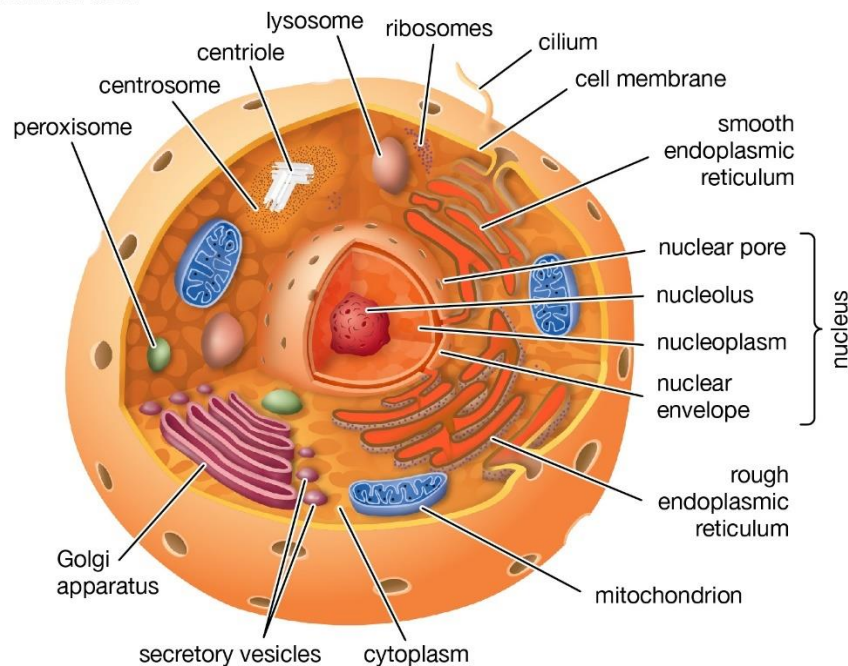
Wir werden uns schwerpunktmäßig mit der Eucyte auseinandersetzen. Denn zum einen gehören unsere Zellen selbst zu dieser Organisationsform der Zelle und zum anderen mögen Bakterien in ihrem Aufbau wesentlich einfacher sein, verfügen aber als Gruppe über eine wesentlich höhere Vielfalt an Stoffwechselsystemen. Bakterien können nahezu alle natürlichen organischen Substanzen abbauen und spielen eine

überragende Rolle im Kohlenstoffkreislauf, Stickstoffkreislauf und Schwefelkreislauf auf der Erde. Die meisten Bakterien benötigen organische Substrate als Energiequelle. Einige Bakteriengruppen gewinnen dagegen ihre Stoffwechselenergie durch Oxidation anorganischer Substrate. So gibt es z. B. schwefeloxidierende Bakterien, können also Schwefelverbindungen wie Schwefelwasserstoff (H_2S) oder Sulfitionen wie SO_3^{2-} als Energiequelle nutzen. Andere Bakterien nutzen Eisenverbindungen oder Stickstoffverbindungen wie Ammonium, Nitrat oder Nitrit als Energiequelle. Andere Bakterien sind in der Lage durch Umwandlung von Lichtenergie ihre Stoffwechselenergie zu gewinnen. Wiederrum andere sind in der Lage ohne Sauerstoff Energiequellen zu nutzen. Einige Arten können in extremen Lebensräumen vorkommen, die kein Tier oder keine Pflanze überstehen würde. Z. B. gibt es Prokaryoten, die in hohen Salzkonzentrationen oder in extremen Temperaturen. Das Wasserstoff-oxidierende Archaeobakterium *Acidianus infernus* wächst zwischen 65 und 96 °C mit einem Wachstumsoptimum bei ca. 90 °C. Der chemolithotrophe Schwefelreduzierer *Pyrodictium occultum* hat sogar ein Temperaturoptimum von 105 °C; unter 82 °C ist kein Wachstum zu beobachten.

Diese Artikelserie kann diesen Extremisten unter den Zellen keinesfalls gerecht werden und ein spezielles Gebiet der Biologie, die sog. Mikrobiologie, widmet sich vornehmlich dem Aufbau, Stoffwechsel und Ökologie der Bakterien und Archaeobakterien. Wir konzentrieren uns hier aber tatsächlich auf die Eucyten, also die Eukaryoten. In ihrem Stoffwechsel sind sie geradezu langweilig, aber durchaus komplex und interessant genug sich näher damit zu befassen. Denn bedenkt: das passiert tagein tagaus in euren eigenen Zellen.

Befassen wir uns erstmal also mit dem Aufbau unserer Zellen (Abb. 2)

Animal cell



© Encyclopædia Britannica, Inc.

Abbildung 2 Zellorganellen tierischer Zellen

Jede Zelle verfügt über Zellplasma, auch Cytoplasma genannt. Das Zellplasma ist die Grundstruktur einer Zelle, die diese ausfüllt. Sie besteht hauptsächlich (zu über 80%) aus Wasser, der Rest besteht überwiegend aus Proteinen, Zuckern, Nukleinsäuren, Fetten etc. Das Cytoplasma ist für den Transport von Stoffen (Moleküle, Enzyme, Nährstoffe) innerhalb der Zelle zuständig. Eine weitere Funktion ist die Lagerung von Wasser und Nährstoffen. Ebenso stellt das Zellplasma Areale zur Verfügung, welche essentiell für bestimmte chemische Vorgänge und biologische Prozesse sind, da diese eine ganz spezifische Umgebung benötigen.

Jede Zelle wird durch eine sog. Zellmembran von ihrer Umgebung abgetrennt. Die Zellmembran besteht aus einer zweilagigen Schicht von Lipiden, also Fettverbindungen (Lipiddoppelschicht) und verschiedenen Membranproteinen, die darin eingelassen sind. Die Zellmembran ist unterschiedlich gut bzw. schlecht durchlässig, sie ist also semipermeabel. Die Zellmembran die Grenzfläche dar, über die ein Stoffaustausch mit der Umgebung stattfindet.

Innerhalb der Zelle befinden sich die sog. Zellorganellen, die jeweils verschiedene Funktionen erfüllen. Jedes Zellorganell ist ebenso abgeschlossen durch Membranen. Wir werden hier die wichtigsten Zellorganellen nennen:

1. Zum einen haben wir den Zellkern, der das Genom beinhaltet. Abgegrenzt wird der Zellkern durch die Kernhülle, die aus einer inneren und einer äußeren Membran besteht. Ein Stoffaustausch findet nur durch die Kernporen mit dem Cytoplasma statt.
2. Ein Zellorganell hat einen ganz komplizierten Namen: das Endoplasmatische Reticulum, auch kurz ER genannt. Es gehört zu dem inneren Membransystem einer Zelle und durchzieht das gesamte Cytoplasma und steht in einem engen Kontakt zur Kernhülle. Beim Endoplasmatischen Retikulum (ER) unterscheidet man zwischen dem rauen ER und dem glatten ER. Am rauen ER sind viele Ribosomen angeheftet, welche Proteine herstellen. Diese werden vom ER weiter verändert und zu einer bestimmten Raumstruktur gefaltet. Am glatten ER hingegen gibt es keine Ribosomen, stattdessen ist diese Struktur für den Abbau von Giften und Medikamenten zuständig oder auch an der Bildung von Lipiden oder Hormonen beteiligt. Weiterhin speichert das ER Calcium, welches eine wichtige Rolle bei der Signalübertragung spielt.
3. Der Golgi-Apparat ist das Zentrum des Stofftransports in der Zelle. Er besteht aus abgeflachten Membransäckchen (Dictyosomen) und Golgi-Vesikeln. Aufgabe des Golgi-Apparats ist es, den Stofftransport in der Zelle mittels Exo- und Endocytose zu koordinieren und Proteine durch die Zelle zu bestimmten Orten zu transportieren.
4. Lysosomen sind Vesikel, in denen sich bestimmte Abbauenzyme befinden. Diese dienen dazu, verschiedene Stoffe, die von der Zelle aufgenommen wurden, zu zersetzen. Dabei entstehen für die Zelle nutzbare Produkte. Eine andere Funktion von Lysosomen besteht darin, beschädigte Zellbestandteile abzubauen, sodass die Bausteine wieder von der Zelle genutzt werden können.
5. Mitochondrien sind Zellorganellen, in denen im Zuge der Zellatmung Energie produziert wird. Daher werden sie auch als Kraftwerke der Zelle bezeichnet. Mitochondrien besitzen zwei Membranen und sowohl eigene DNA als auch Ribosomen. Sie waren ursprünglich Bakterien, die von der Ur-Eucyte aufgenommen, aber nicht zersetzt wurden. man spricht von der sog. Endosymbiontentheorie.

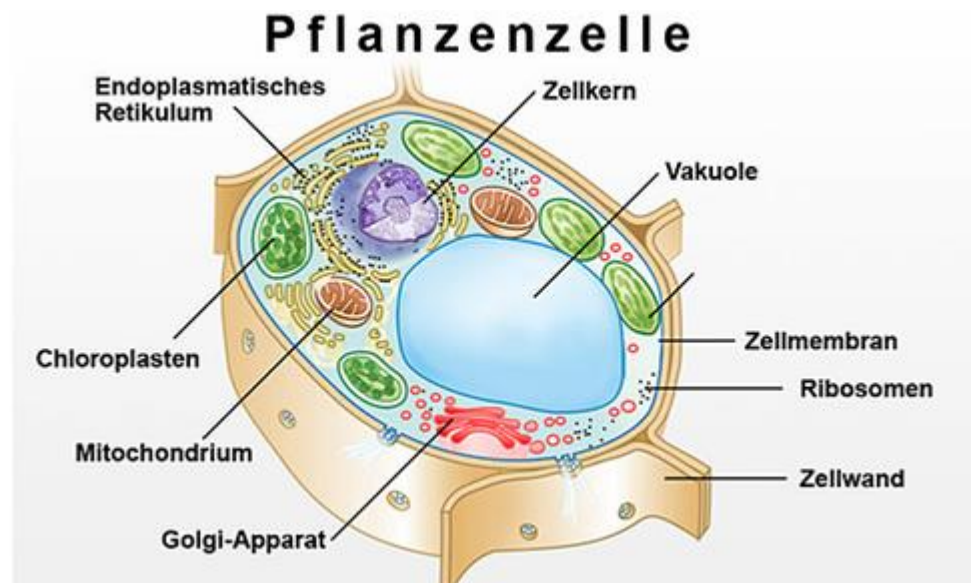


Abbildung 3 Pflanzenzelle

1. Pflanzliche Zellen haben zusätzliche Organellen, die in tierischen fehlen (Abb. 3). Hierzu zählen die Chloroplasten. Sie sind, genau wie Mitochondrien, von zwei Membranen umgeben und besitzen eine eigene DNA, waren also auch ursprünglich Bakterien, die von der Ur-Pflanzenzelle durch Endosymbiose aufgenommen wurden. Chloroplasten besitzen im Inneren viele flache Membransäckchen (Thylakoide), die in Stapeln (Grana) angeordnet sind. Hier ist das Chlorophyll eingelagert. An den Chloroplasten findet die Photosynthese statt.
2. Ein weiteres Zellorganell der Pflanzen sind die Vakuolen. Sie nehmen in pflanzlichen Zellen oft den größten Teil der Zelle ein. Sie sind mit Zellsaft gefüllt und dienen zum einen der Speicherung von Stoffen und geben zum anderen der Zelle Stabilität.
3. Pflanzenzellen verfügen zusätzlich über eine Zellwand. Die Zellwand liegt außerhalb der Zellmembran, die ihrerseits das Zellinnere enthält. Sie wird als Abfallsprodukt lebender Zellen gebildet. Die Zellwand bietet Struktur und Schutz und wirkt zudem als Filter. Hauptbestandteil der Zellwand bei Pflanzen ist Cellulose. Aber auch Pilze und Bakterien haben eine Zellwand. Die Zellwand der meisten Bakterien besteht aus dem Peptidoglykan Murein, bei Pilzen aus Chitin.

Zellen verfügen zusätzlich im Cytoplasma das Cytoskelett. Das Cytoskelett ist ein Zellorganell in eukaryotischen Zellen. Die Prokaryoten verfügen über kein Cytoskelett im eigentlichen Sinne, enthalten aber trotzdem ähnliche Proteinstrukturen. Das Cytoskelett beschreibt dabei ein Netzwerk im Cytoplasma, das aus mehreren Proteinen und länglichen Filamenten aufgebaut ist. Seine drei Hauptbestandteile in den Eukaryoten sind die Mikrotubuli, die Mikrofilamente und die Intermediärfilamente. Seine Aufgaben bestehen in der mechanischen Stabilisierung, im Stofftransport und in der aktiven Bewegung der Zelle.

Kapitel 2: Atome, Periodensystem und Atombindung

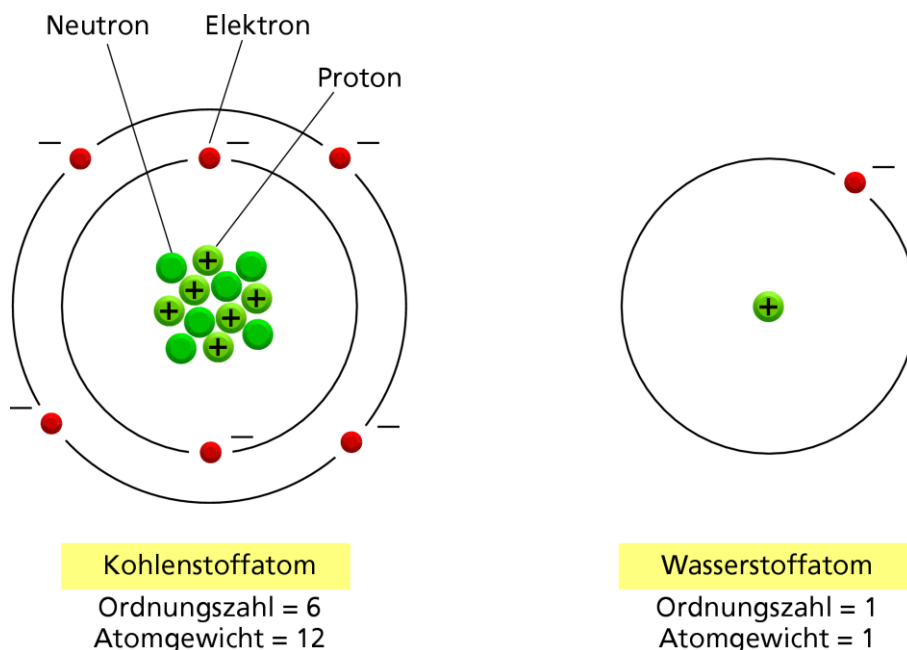
Im ersten Kapitel befassten wir uns mit dem grundlegenden Aufbau der Zelle.

In diesem Teil lernen wir, wie Atome aufgebaut sind und sie sich miteinander verbinden, etwas über das Periodensystem der Elemente und zum Ende wird es ein wenig nass, denn es geht ums Wasser.

Atome – Grundaufbau

Doch fangen wir ganz klein an. Wir haben gelernt, dass die Zelle die kleinste Einheit des Lebens ist. Doch in der Welt der Materie geht es natürlich kleiner. Zellen bestehen aus verschiedenen Verbindungen: Wasser, Proteine, Fette, Zucker, Vitamine usw. usf. Hierbei sprechen wir von Molekülen. Moleküle sind wiederum aus Atomen zusammengesetzt.

Das Atom ist jeweils der kleinste Teil der Materie eines Elementes. Eine lange Zeit nahm man an, dass Atome der kleinste überhaupt existierende Teil der Materie sind. Atom heißt nämlich: unteilbar. Heute wissen wir natürlich, dass Atome ihrerseits aus subatomaren Teilchen bestehen. Die für uns bekanntesten subatomaren Teilchen und die, die für unsere Videoreihe völlig ausreichen, sind die Protonen und Neutronen, sowie die Elektronen (Abb. 4).



© 2012 Wiley-VCH, Weinheim
Alberts – Lehrbuch der Molekularen Zellbiologie
ISBN: 978-3-527-32824-6 Fig. 02-002

Abbildung 4 Schematische Darstellung eines Kohlenstoffatoms und eines Wasserstoffatoms. Der Kern jedes Atoms außer Wasserstoff besteht aus positiv geladenen Protonen und elektrisch neutralen Neutronen. Die Zahl der Elektronen eines Atoms ist gleich der Protonenzahl, sodass das Atom keine Nettoladung trägt. Die Neutronen, Protonen und Elektronen sind in Wirklichkeit winzig im Vergleich zum Gesamtatom; ihre Größen sind hier stark übertrieben.

Also: Schauen wir uns die Abb. 4 genauer an: Die Protonen sind die positive Ladung des Atoms und befinden sich im Atomkern. Ebenfalls im Atomkern befinden sich die Neutronen, die, wie ihr Name schon sagt, keine Ladung haben. Auf der Atomhülle befinden sich die negativ geladenen Elektronen. Der Atomkern bildet die größte Masse eines Atoms, nimmt aber den kleinsten Platz im Atom ein. Die Atomhülle mit ihren Elektronen nimmt hingegen das größte Volumen des Atoms an, wiegt aber nur einen verschwindend kleinen Bruchteil des Atomkerns. Die jeweiligen Elektronen bewegen sich auf bestimmten Bahnen, die man als Orbitale oder in der vereinfachten Form Schalen bezeichnet. Die konzentrischen schwarzen Kreise repräsentieren stark schematisiert die „Orbitale“ (also die Verteilungen der Elektronen in der Atomhülle) der Elektronen.

Dieser Grundaufbau bei Atomen ist gleich: der Unterschied liegt in der Anzahl dieser subatomaren Teilchen. Das leichteste und kleinste Atom ist der Wasserstoff. Es hat im Kern ein Proton und in der Hülle ein Elektron. Würde man ein zweites Proton und ein zweites Elektron hinzufügen erhält man Helium, fügt man ein drittes Proton und ein drittes Elektron hinzu, erhält man Lithium usw. usf. Wie viele Elektronen sich in der Hülle befinden, hängt von der Anzahl der Protonen im Kern ab. Es befinden sich nämlich genauso viele negativ geladene Elektronen in der Atomhülle wie positiv geladene Protonen im Atomkern. Das gibt die sogenannte Ordnungszahl an. Sie ist beim Wasserstoffatom 1, weil dieser je 1 Proton und 1 Elektron hat. Dadurch ist ein Atom insgesamt ungeladen. Je mehr Elektronen und Protonen hinzugefügt werden, desto schwerer wird das Atom und es entsteht ein anderes Element. Dies zeigt sich am Beispiel des Kohlenstoffatoms. Es hat die Ordnungszahl 6, hat also 6 Protonen und 6 Elektronen. Neben der Ordnungszahl gibt es das Atomgewicht. Elektronen sind sehr klein und haben vielleicht 0,0005% der Atommasse, sind also vernachlässigbar. Die Hauptmasse des Atoms befindet sich im Atomkern, dabei sind Protonen und Neutronen etwa gleich schwer. Kohlenstoff hat ein Atomgewicht von 12. Es hat 6 Protonen im Atomkern und 6 Neutronen. Beide zusammen ergeben das Atomgewicht.

Die Elektronen bewegen sich auf sogenannten Orbitalen, die oft in einem Schalenmodell vereinfacht dargestellt werden. Die innerste Schale trägt den Kennbuchstaben K. Jede weitere Schale ist dann mit fortlaufenden Buchstaben des Alphabets (L, M, N, usw.) betitelt. In jeder Schale gibt es eine maximale Anzahl an Elektronen. Für die K-Schale sind das zwei Elektronen. In die L-Schale passen höchstens acht Elektronen bis sie gefüllt ist und für die M-Schale beträgt die Maximalzahl 18 Elektronen, in der N-Schale haben 32 und in der O-Schale 50 Elektronen Platz.

Periodensystem der Elemente (PSE)

Auf diesem Prinzip beruht übrigens das Periodensystem der Elemente (Abb. 5). Dies ist das Ordnungssystem, die verschiedenen Elemente zu gruppieren. Der russische Chemiker Dimitri Mendelejew und der deutsche Chemiker Lothar Meyer haben dies unabhängig voneinander fast zeitgleich in den 1860er Jahren erkannt. Beide haben ihre Erkenntnisse Jahr 1869 in einer wissenschaftlichen Zeitschrift veröffentlicht.

I																	VIII				
1,01 H 1																	4,00 He 2				
6,94 Li 3	9,01 Be 4															10,81 B 5	12,01 C 6	14,01 N 7	16,00 O 8	19,00 F 9	20,18 Ne 10
22,99 Na 11	24,31 Mg 12	III a	IV a	V a	VI a	VII a	VIII a		I a	II a	13 Al	14 Si	15 P	16 S	17 Cl	18 Ar					
39,10 K 19	40,08 Ca 20	44,96 Sc 21	47,87 Ti 22	50,94 V 23	52,00 Cr 24	54,94 Mn 25	55,85 Fe 26	58,93 Co 27	58,69 Ni 28	63,55 Cu 29	65,39 Zn 30	69,72 Ga 31	72,61 Ge 32	74,92 As 33	78,96 Se 34	79,90 Br 35	83,8 Kr 36				
85,47 Rb 37	87,62 Sr 38	88,91 Y 39	91,22 Zr 40	92,91 Nb 41	95,94 Mo 42	97,91 Tc 43	101,0 Ru 44	102,9 Rh 45	106,4 Pd 46	107,9 Ag 47	112,4 Cd 48	114,8 In 49	118,7 Sn 50	121,8 Sb 51	127,6 Te 52	126,9 I 53	131,3 Xe 54				
132,9 Cs 55	137,3 Ba 56	175,0 Lu 71	178,5 Hf 72	180,9 Ta 73	183,8 W 74	186,2 Re 75	190,2 Os 76	192,2 Ir 77	195,1 Pt 78	197,0 Au 79	200,6 Hg 80	204,4 Tl 81	207,2 Pb 82	209,0 Bi 83	209,0 Po 84	210,0 At 85	222,0 Rn 86				
223,0 Fr 87	226,0 Ra 88	262,0 Lr 103	261,1 Rf 104	262,1 Db 105	266,1 Sg 106	264,1 Bh 107	269,1 Hs 108	268,1 Mt 109	273,1 Ds 110	272,1 Rg 111											

Atommasse in u
(molare Masse)

26,98

Al

13

Elementsymbol

Ordnungszahl

Wasserstoff

radioaktiv

Erdalkalimetalle

Metalle

Halbmetalle

Edelgase

Nichtmetalle

Alkalimetalle

Abbildung 5 Periodensystem der Elemente

Grundsätzlich kann man den Aufbau des Periodensystems wie folgt beschreiben:

Alle Elemente werden durch ein Elementsymbol beschrieben. Wasserstoff wird mit H abgekürzt, Sauerstoff mit O, Kohlenstoff mit C und Stickstoff mit N. In Abb. 5 wird das Element Aluminium, abgekürzt durch Al hervorgehoben. An diesem werden wir den Rest des Periodensystems erklären.

[illegible]

Aluminium hat eine Atommasse von 26,96 u, also gerundet etwa 27. Wenn Aluminium schon 13 Protonen hat, beträgt die Differenz zu 27 u, die Anzahl der Neutronen. $27 - 13 = 14$ Neutronen.

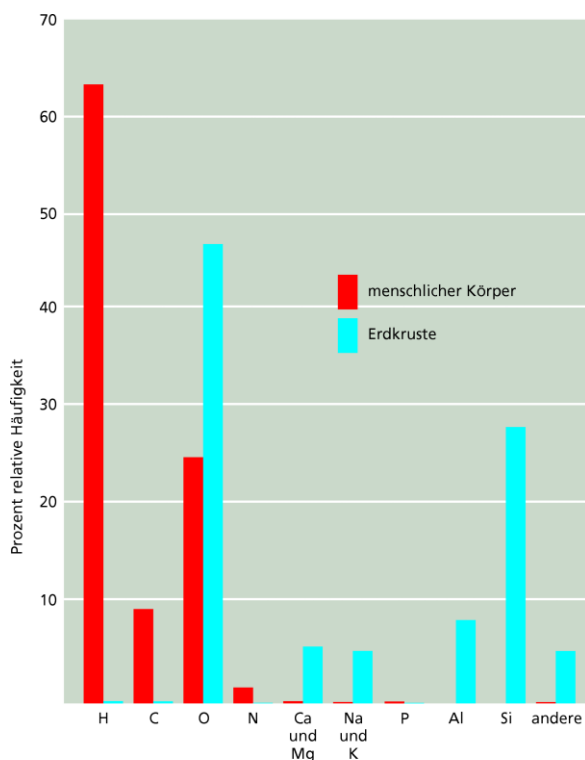
Das Periodensystem ist angeordnet in Zeilen und Spalten.

Die Zeilen sind die Perioden und sagen aus, wie viele Elektronenschalen das Element hat. Aluminium ist in der dritten Periode, hat also 3 Elektronenschalen, die letzte Schale ist die M-Schale.

Die Spalten stellen die Gruppen da. Es gibt 8 Hauptgruppen, dargestellt durch die römischen Zahlen. Zwischen Hauptgruppe II und III befinden sich die Nebengruppen, die wir erstmal ignorieren können. Gruppen sind Elemente mit ähnlichen Eigenschaften. Die Gruppe gibt an, wie viele Elektronen sich auf der äußersten Schale befinden. Aluminium ist in Hauptgruppe III. Auf seiner äußersten M-Schale befinden sich also 3 Elektronen. Das macht Sinn: Aluminium hat 13 Elektronen (siehe Ordnungszahl) und ist in Periode drei, hat also die K, L und M-Schale. Die K-Schale kann maximal 2 Elektronen aufnehmen, die L-Schale maximal 8. Damit hätten wir schon 10 Elektronen. Fehlen noch 3 Elektronen, um auf die 13 zu kommen. Die befinden sich nun auf der dritten Schale, also der M-Schale.

Übrigens hatte ich erwähnt, dass die M-Schale bis zu 18 Elektronen aufnehmen kann, die N-Schale (Periode 4) 31 und die O-Schale (Periode 5) sogar 50. Doch bei den 8 Hauptgruppen hat die äußerste Schale maximal 8 Elektronen (also Valenz- oder Außenelektronen).

Die verschiedene Farbgebung zeigt des Weiteren den Unterschied zwischen Metallen, Nichtmetallen etc. an. Für die Chemie des Lebens sind von den 111 Elementen tatsächlich ungleichverteilt. Abb. 6 zeigt den Anteil der wichtigsten Elemente im menschlichen Körper (rot) und auf der Erdkruste (blau). Es fällt auf: der menschliche Körper besteht zu 60% aus Wasserstoffatomen.



© 2012 Wiley-VCH, Weinheim
Alberts - Lehrbuch der Molekularen Zellbiologie
ISBN: 978-3-527-32824-6 Fig. 02-004

Abbildung 6 Die Häufigkeit jedes Elements ist als Prozent der Gesamtzahl der Atome in einer biologischen oder geologischen Probe angegeben, einschließlich Wasser. Demnach sind mehr als 60 % der Atome in einem lebenden Organismus Wasserstoffatome.

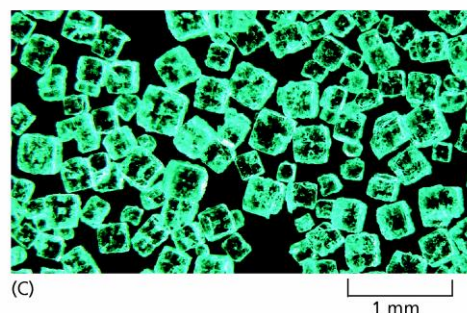
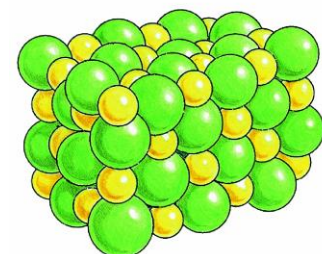
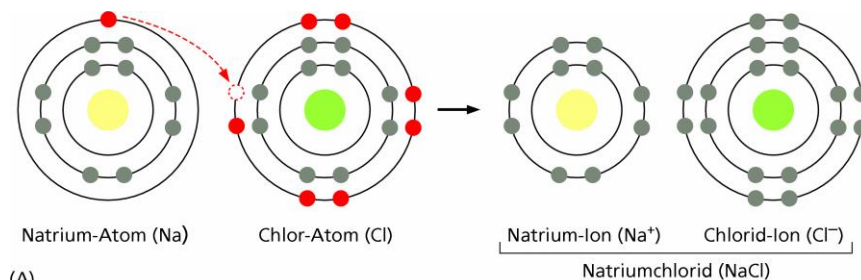
Ionen und Ionenbindung

Atome haben den Drang dazu, ihre äußerste Schale voll zu kriegen, das wäre bei den Hauptgruppen 8 Elektronen auf der Außenschale (mit Ausnahme der ersten Schale, die maximal zwei Elektronen aufnehmen kann). Dies bezeichnet man als Edelgaszustand. Wenn man sich das PSE (Abb. 5) anschaut, sind die Edelgase Helium, Neon, Argon, Krypton und Xenon in der Hauptgruppe VIII. Damit ist ihre äußerste Stelle mit 8 Elektronen besetzt – diese Atome sind, wenn man es so will, glücklich. Doch was machen die anderen Elemente, die ebenfalls nach diesem Edelgaszustand streben? Es gibt mehrere Möglichkeiten:

Eine Möglichkeit wäre, überzählige Außenelektronen abzugeben oder fehlende aufzunehmen. Das ist bei Ionen der Fall. Denn geben sie Elektronen ab oder nehmen welche auf, entspricht ihre Zahl an Elektronen nicht mehr der Zahl der Protonen. Damit sind sie nicht mehr elektrisch neutral geladen, sondern positiv geladen (wenn sie zu wenig Elektronen haben) oder negativ (wenn sie mehr Elektronen als Protonen haben). Positiv geladene Ionen heißen Kationen, negativ geladene Ionen Anionen. Die Ladung eines Ions steht dann oben rechts am Atomsymbol.

Wenn man das Periodensystem soweit verstanden hat, kann man sich – zumindest was die Hauptgruppen angeht – gut vorstellen, welche Elemente eher Elektronen abgeben und welche aufnehmen.

Für die Elemente der ersten und zweiten Hauptgruppe ist es z. B. leichter ihre eine bzw. zwei Elektronen auf der Außenschale abzugeben. Elemente in den Hauptgruppen VI und VII neigen dann eher zusätzliche Elektronen aufzunehmen. Das kann man sehr gut an zwei Elementen zeigen: Natrium und Chlor (Abb. 7).



© 2012 Wiley-VCH, Weinheim
Alberts - Lehrbuch der Molekularen Zellbiologie
ISBN: 978-3-527-32824-6 Fig. 02-008

Abbildung 7 Natrium und Chlor-Ionen. A) Ein Natriumatom (Na) reagiert mit einem Chloratom (Cl). Die Elektronen jedes Atoms sind in ihren verschiedenen Energieniveaus dargestellt. Die Elektronen in den chemisch reaktiven (unvollständig gefüllten) Schalen sind rot gezeichnet. Die Reaktion findet durch die Übertragung eines einzelnen Elektrons von Natrium auf Chlor statt.

Dabei entstehen zwei elektrisch geladene Atome, oder Ionen, mit jeweils vollständigem Elektronensatz in der äußeren Schale. Die beiden Ionen mit entgegengesetzten Ladungen werden durch elektrostatische Anziehung zusammengehalten. (B) Das Produkt der Reaktion zwischen Natrium und Chlor ist das kristalline Natriumchlorid. Es besteht aus Natrium- und Chloridionen, die in einer regelmäßigen Anordnung dicht gepackt sind; dabei gleichen sich die Ladungen genau aus. (C) Farbphotographie von Natriumchloridkristallen.

Natrium befindet sich in der ersten Hauptgruppe und Periode 3. Es hat also 3 Elektronenschalen ein Außenelektron auf der dritten Schale. Dieses gibt er dann ab, so dass die dritte Schale leer wird und er mit der zweiten Schale den Edelgaszustand erreicht. Nun hat er aber ein Elektron zu wenig und wird dadurch positiv, also Na^+ . Chlor ist in der siebten Hauptgruppe und Periode drei. Es hat also 3 Elektronenschalen und auf der äußersten befinden sich 7 Elektronen. Um den Edelgaszustand zu erreichen, fehlt nur ein Elektron auf der äußersten Schale. Nimmt es ein zusätzliches Elektron auf, hat es den Edelgaszustand erreicht und ist negativ geladen: Cl^- .

Die Fähigkeit zusätzliche Elektronen aufzunehmen, bezeichnet man als Elektronegativität.

Positive und negative Ladungen ziehen sich gegenseitig an. Wenn dies geschieht entstehen Ionenbindungen. Das ist vor allem bei Salzen der Fall. Wenn sich Natrium und Chlor gegenseitig anziehen, entsteht Natriumchlorid. Das ist das Salz welches wir in unserem Essen benutzen: übrigens macht es da keinen Unterschied ob wir Himalaya-Salz oder Meersalz nutzen. Wie andere rohe (unraffinierte) Steinsalze besteht es zu 97 bis 98 Prozent aus Natriumchlorid und einem geringen Anteil von anderen Mineralen wie Gips und Kaliumchlorid. Seine rötliche Färbung verdankt es den Eisenionen eines geringfügigen Anteils an Eisenoxidverunreinigungen. Einen gesundheitlichen Vorteil bietet das Himalayasalz nicht.

Wasserstoff z. B. gibt auch häufig sein Elektron ab und ist dadurch positiv geladen. Da ein positiv geladenes Wasserstoff-Atom (H^+) nur aus einem Proton besteht, werden sie auch oft nur Protonen genannt.

Atombindung

Eine weitere Möglichkeit den Edelgaszustand zu erreichen ist die Elektronenpaarbindung, auch Atombindung oder kovalente Bindung genannt. Dies kann am Beispiel des Wassers verdeutlicht werden.

Wasser besteht aus einem Sauerstoffatom und zwei Wasserstoffatomen. Seine Verbindung wird entsprechend H_2O genannt. H steht für das Wasserstoff, weil es zwei davon gibt, ist die Zahl zwei unten eingefügt, O steht für das Sauerstoff. Aber warum diese Verbindung?

Nun schauen wir ins Periodensystem: H steht in der Hauptgruppe 1 und hat 1 Außenelektron. Auf seine Schale passen maximal zwei Elektronen, dann hätte es den Edelgaszustand, ähnlich wie Helium erreicht. Sauerstoff steht in Hauptgruppe VI. Es hat also auf seiner Außenschale sechs Elektronen. Um den Edelgaszustand zu

erreichen, tun sich die zwei Wasserstoffatome mit einem Sauerstoffatom zusammen (Abb. 8). Wasserstoff und Sauerstoff teilen sich entsprechend ein Elektronenpaar, so dass alle drei Atome jeweils eine volle Schale haben. Der Sauerstoff stellt zwei seiner sechs Außenelektronen den beiden Wasserstoffatomen zu Verfügung. Die vier verbliebenen Außenelektronen des Sauerstoffs bilden zwei freie Elektronenpaare. Im Gegensatz stellt das Wasserstoffatom dem Sauerstoffatom sein Elektron zu Verfügung. Mit jedem Wasserstoff geht der Sauerstoff jeweils eine Einfachbindung ein, da sie sich ein Elektronenpaar teilen.

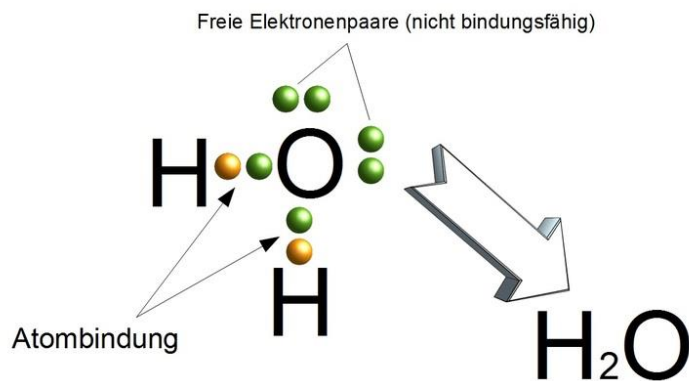


Abbildung 8 Atombindung des Wassers

Neben einer Einfachbindung gibt es auch Doppel- und Dreifachbindungen.

Eine Doppelbindung kommt z. B. beim CO_2 vor, dieses besteht aus einem Kohlenstoffatom, abgekürzt durch das C und zwei Sauerstoffatomen. Sauerstoff hat wie gesagt 6 Außenelektronen, hätte gerne zwei weitere für den Edelgaszustand. Kohlenstoff befindet sich in der IV. Hauptgruppe, hat also vier Außenelektronen. Um den Edelgaszustand zu erreichen, also 8 Elektronen auf der Außenschale zu besitzen, teilt das Kohlenstoffatom seine 4 Außenelektronen mit zwei Sauerstoffatomen. Mit jedem Sauerstoffatom teilt es dann zwei Elektronenpaare. Ein gemeinsames Elektronenpaar kann statt durch zwei Punkte auch durch einen Strich dargestellt werden, wie Abb. 9 zeigt.

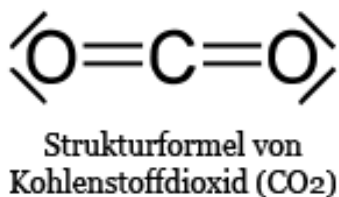


Abbildung 9 Strukturformel des CO_2

Isotope

Isotope sind Elemente mit gleicher Ordnungszahl, aber unterschiedlicher Massenzahl. D. h. zwei Isotope haben die gleiche Anzahl an Protonen (und Elektronen), aber die Zahl der Neutronen im Kern ist unterschiedlich. Dies kann am Beispiel des Wasserstoffs (H) verdeutlicht werden, von dem es drei Isotope gibt (Abb. 10):

Es gibt das „gewöhnliche“ Wasserstoff auch Protium genannt, das nur ein Proton im Atomkern hat. Es stellt übrigens 99,99% aller Wasserstoffisotope dar. Deuterium hat im Atomkern ein zusätzliches Neutron, Tritium zwei zusätzliche Neutronen. Deuterium und Tritium kommen in sehr geringen Mengen auf der Erde vor.



Abbildung 10 Isotope des Wasserstoffes

Wasser und pH-Wert

Das für uns wichtigste Molekül ist Wasser. Wasser zählt nicht nur für den Menschen zu den wichtigsten Stoffen auf der Erde, sondern das Leben ist untrennbar mit der Existenz von Wasser verbunden. Unser Körper besteht zu 70% aus Wasser. Klingt viel ist es aber nicht, wenn man bedenkt, dass Quallen zu 98-99% aus Wasser bestehen.

Mehr als die Hälfte des Wassers, etwa 55%, befindet sich innerhalb unserer Zellen. Der Rest des Wassers ist extrazellulär und findet sich im Blutplasma (8%), der Lympheflüssigkeit (22%), sowie im Bindegewebe, Knochen und Knorpel (15%).

Wasser wird auch als Lösungsmittel für biochemische Reaktionen gebraucht, sei es beim Stofftransport durch die Zellmembranen, der Erhaltung unserer Körpertemperatur oder als Lösungsmittel für unsere Ausscheidungen.

Chemisch betrachtet besteht Wasser, wie wir nun wissen, aus zwei Wasserstoffatomen und einem Sauerstoffatom. Diese Verbindung, obwohl sie so einfach ist, hat entscheidende physikalische Eigenschaften, die das Leben ermöglichen.

Wasser ist ein polares Molekül. Polarität bezeichnet in der Chemie eine durch Ladungsverschiebung in Atomgruppen entstandene Bildung von getrennten

Ladungsschwerpunkten, die bewirken, dass eine Atomgruppe nicht mehr elektrisch neutral ist.

Schauen wir uns das Wassermolekül näher an (Abb. 8). Wir sehen, dass die Wasserstoffatome am Sauerstoffatom angewinkelt sind. Der Bindungswinkel zwischen den drei Atomen beträgt ungefähr $104\text{--}105^\circ$. Aber warum diese Winkelung? Das hat mit den zwei freien Elektronenpaaren des Sauerstoffs zu tun, die nicht in die Bindung mit dem Wasserstoff eingehen. Die brauchen ja Platz und entsprechend der Winkel. Das hat aber Konsequenzen für die Ladung des Wassermoleküls.

Beim Wassermolekül ist der mittlere Teil, also in der Nähe des Sauerstoffatoms mit den zwei freien Elektronenpaaren, negativer geladen, nach außen hin, in Richtung Wasserstoff, positiver. Die negativ geladenen Elektronen der Wasserstoffatome neigen also in Richtung des Sauerstoffatoms. Wasser ist also polar. Das hat mit der sog. Elektronegativität zu tun. Elektronegativität ist ein Maß für die Fähigkeit, von beiden Bindungspartnern geteilten Elektronen wie beim Tauziehen auf seine Seite zu holen. Und Sauerstoff hat eine höhere Elektronegativität als Wasserstoff, ist also beim Elektronentauziehen gegenüber dem Wasserstoff auf der Gewinnerseite. Diese Eigenschaft des Wassers macht ihn zu einem ganz besonderen Molekül. CO_2 hingegen ist unpolar. Seine Bindungswinkel liegen bei 180° (Abb. 11).

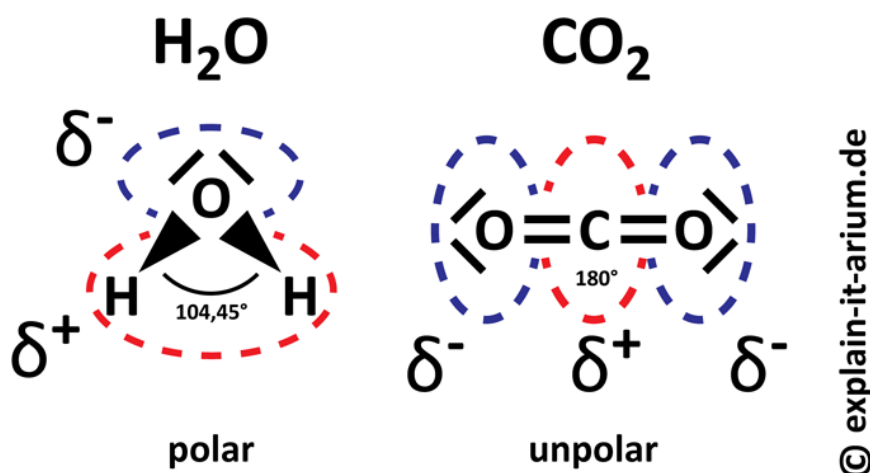


Abbildung 11 Polarität des Wassers.

Zwischen den Wassermolekülen mit ihren teilweise positiven und teilweise negativen Ladungsenden treten Wechselwirkungen auf, die man als Wasserstoffbrückenbindungen bezeichnet werden. Das heißt eigentlich nur, dass ein Wassermolekül ein anderes Wassermolekül anzieht. Der teilweise positiv geladene Wasserstoffteil des einen Wassermoleküls zieht den teilweise negativ geladenen Sauerstoffatom des anderen Wassermoleküls an (Abb. 12). Solche Wasserstoffbrückenbindungen kommen auch bei anderen Molekülen vor, die für die Zellbiologie sehr wichtig sind, so bei der DNA.

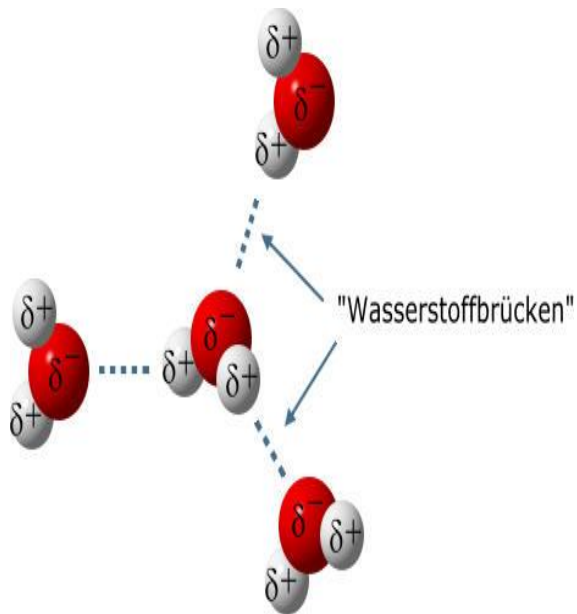


Abbildung 12 Wasserstoffbrückenbindung

Diese Wasserstoffbrückenbindungen spielen in der Biochemie eine sehr wichtige Rolle, da sie die Reaktionsfähigkeit der Moleküle stark beeinflussen. Für das Wasser hat es besonders wichtige Konsequenzen:

Wasser hat eine besondere Eigenschaft, die es von fast allen anderen Flüssigkeiten unterscheidet. Es hat bei 4 °C sein kleinstes Volumen und damit seine größte Dichte. Bei 0 °C wird Wasser fest, also zu Eis. Bei anderen Stoffen haben feste Aggregatzustände eigentlich eine höhere Dichte. Nicht so beim Wasser. Eis hat eine geringere Dichte als 4°C warmes flüssiges Wasser. Deswegen schwimmt in einem gefrorenen Gewässer Eis oben. Die Eisschicht, die sich an der Oberfläche kalter Gewässer bildet, isoliert die flüssige Phase von einer noch kälteren Luft und bewahrt somit das Leben unter dem Eis vor dem Erfrieren! Dieses nicht normale thermische Verhalten von Wasser wird in der Physik als Anomalie des Wassers bezeichnet (Abb. 13)

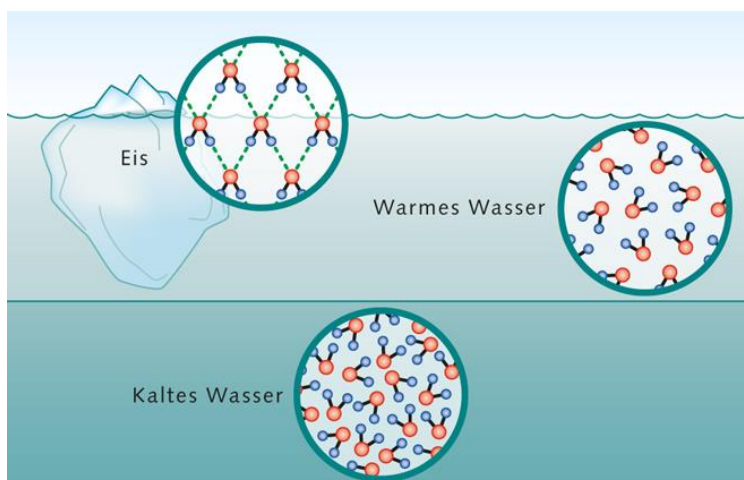


Abbildung 13 Dichte des Wassers. Wasser mit 4°C hat die höchste Dichte

Wasser besitzt eine hohe spezifische Wärmekapazität. Das ist die Energie, die nötig ist, um die Temperatur von einem kg Wasser um ein Grad Celsius zu erhöhen. Eine hohe Wärmekapazität bedeutet, dass es nicht so leicht ist, Wasser zu erhitzen. Außerdem hat Wasser eine hohe Verdampfungsenthalpie. Die Verdampfungsenthalpie ist die Energie, die aufgebracht werden muss, um eine gegebene Menge einer Flüssigkeit bei gegebener Temperatur zu verdampfen (Abb. 14). Aufgrund der hohen Verdampfungsenthalpie und Wärmekapazität des Wassers können Meere und Seen eine beträchtliche Menge Energie aufnehmen und abgeben, ohne dass sich die Temperaturen drastisch ändern würden. Dieses Phänomen hilft auch bei der Kontrolle der Körpertemperatur der Organismen.

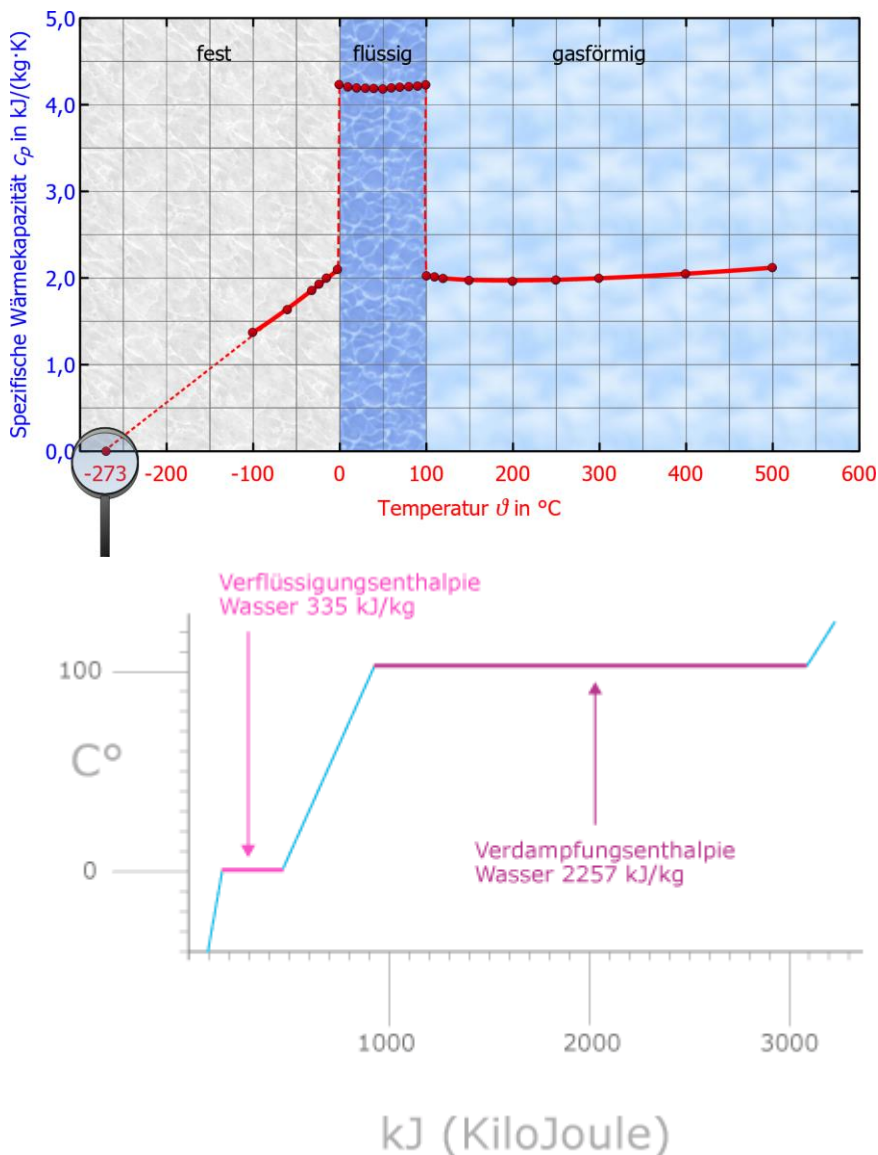


Abbildung 14 Wärmekapazität und Verdampfungsenthalpie von Wasser

Aber die wichtigste Eigenschaft des Wassers: Es ist ein optimales Lösungsmittel. Im Wasser können sich problemlos andere polare Substanzen oder Ionen auflösen. Deswegen löst sich Salz im Wasser, denn Salze sind Ionenverbindungen. Und deswegen können wir Bier und Cola trinken: denn Alkohol und Zucker haben ebenso polare Eigenschaften. Unpolare Stoffe, z. B. Fette und Öle, sind nicht polar und lösen

sich nicht im Wasser. Polare Moleküle werden aufgrund ihrer Fähigkeit sich im Wasser zu lösen als hydrophil bezeichnet, unpolare Moleküle, die sich nicht im Wasser lösen, als hydrophob. Haben Moleküle sowohl hydrophile wie hydrophobe Eigenschaften oder Enden, sind sie amphipatisch. Hört sich alles irgendwie nach psychologischen Störungen an ...

In wässrigen Lösungen – besonders in biologischen Systemen – ist die Konzentration der Wasserstoffionen, also der Protonen (H^+) von Bedeutung. In wässrigen Lösungen neigt das Wassermolekül (H_2O) in äußerst geringem Maße der sog. elektrolytischen Dissoziation. Unter Dissoziation versteht man in der Chemie den angeregten oder den selbsttätig ablaufenden Vorgang der Zerlegung eines Moleküls in zwei oder mehrere einfachere Moleküle, Atome oder Ionen. Das Wassermolekül dissoziiert in ein Wasserstoff-Ion (H^+) und ein Hydroxid-Ion (OH^-) $\rightarrow H_2O = H^+ + OH^-$ (Abb. 15)

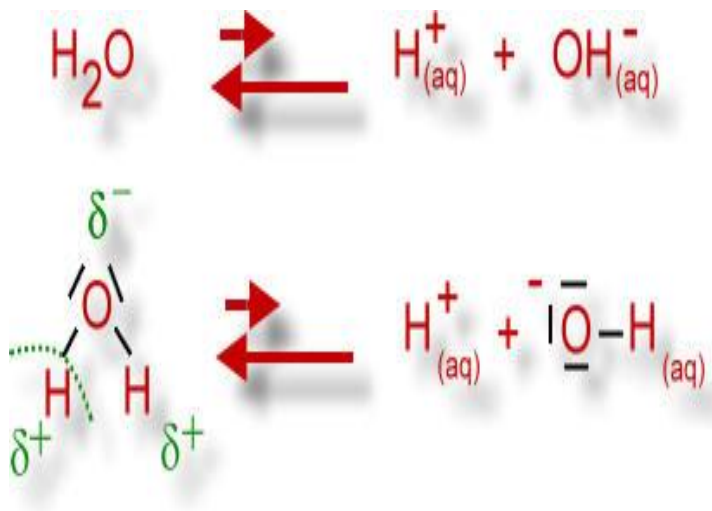


Abbildung 15 Dissoziation von Wasser

Diese Ionen haben bestimmte Eigenschaften: Die Protonen (H^+) wirken als starke Säure, die Hydroxid-Ionen (OH^-) als starke Base. In reinem Wasser liegen diese Ionen bei 25 °C in gleicher Konzentration vor, d.h. im Verhältnis 1:1. Der pH-Wert hat den Wert 7, Wasser ist also neutral. Wenn sich dieses Verhältnis verändert, ändert sich auch der pH-Wert: mehr Protonen (H^+) und weniger Hydroxid-Ionen (OH^-) $\Rightarrow$ Lösung wird saurer $\Rightarrow$ pH-Wert sinkt. Weniger Protonen (H^+) und mehr Hydroxid-Ionen (OH^-) $\Rightarrow$ Lösung wird basischer/alkalischer $\Rightarrow$ pH-Wert steigt (Abb. 16).

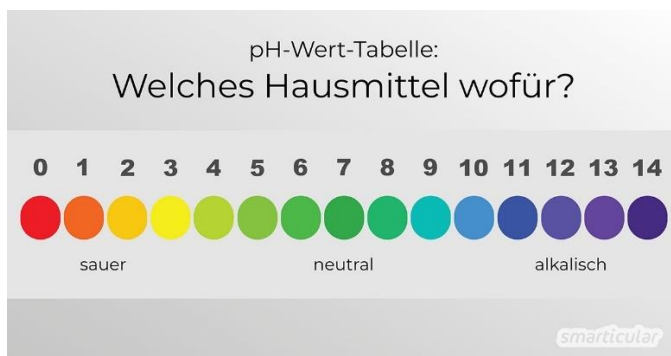


Abbildung 16 pH-Werte

Biologische Systeme sind damit beschäftigt ihren pH-Wert aufrechtzuerhalten. In unserem Blut finden alle lebenswichtigen Stoffwechselvorgänge in einem sehr begrenzten pH-Bereich von 7,35-7,45 statt. Wird dieser Bereich über- oder unterschritten, kommt es zu ernststen Problemen, die im schlimmsten Fall zum Tode führen können. Sinkt der pH-Wert im Blut um mehr als 0,2 Einheiten, das Blut wird also „saurer“, spricht man von Azidose. Ein entsprechender Anstieg führt zur Alkalose (Abb. 17). Wie der pH-Wert im Blut konstant gehalten wird, wird ein anderes Mal behandelt.

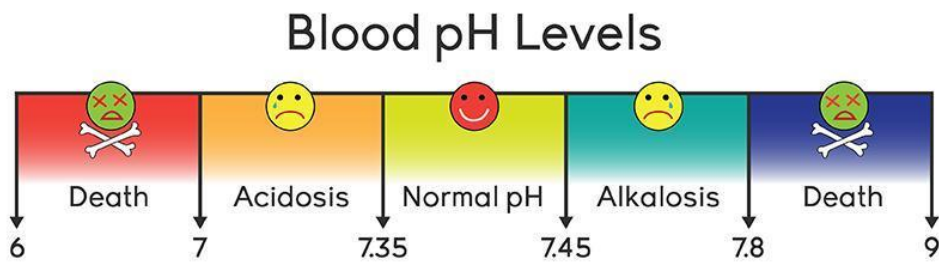


Abbildung 17 pH-Werte des Blutes

Kapitel 3: Proteine, Zucker und Fette

Wir sind organisch!

Nachdem wir nun so einiges über das Wasser erfahren haben, wenden wir uns größeren Molekülen zu. Abgesehen vom Wasser sind die meisten Moleküle in unserem Körper organische Verbindungen. Aber was heißt organische Verbindung überhaupt? Das sind Verbindungen, die Kohlenstoffatome enthalten. Und tatsächlich: alle größeren Moleküle und Naturstoffe haben das Element Kohlenstoff: ob Zucker, ob Eiweiße oder Fette. Doch warum ausgerechnet Kohlenstoff?

Das hat viel mit den Eigenschaften dieses Elements zu tun, denn Kohlenstoff kann wie kein anderes Element vielfältige und vor allem stabile Verbindungen eingehen. Besonderheiten des Kohlenstoffs sind es, Ketten und Ringe mit sich selbst und anderen Elementen sowie Doppel- und Dreifachbindungen zu bilden. Aufgrund seiner mittelstarken Elektronegativität hat er ein gutes Bindungsvermögen sowohl zu elektropositiveren als auch zu elektronegativeren Elementen. Kohlenstoff ist also das Element, dass am besten mit den anderen Elementen zurechtkommt. Für die biochemischen Vorgänge in Lebewesen sind Kohlenstoffverbindungen mit Wasserstoff (H), Stickstoff (N), Sauerstoff (O) und Schwefel (S) am bedeutendsten. Die letzten drei Elemente sind bei organischen Molekülen Hauptbestandteile der funktionellen Gruppen. In der Chemie versteht man unter funktionellen Gruppen Atomgruppen in organischen Verbindungen, die die Stoffeigenschaften und das Reaktionsverhalten bestimmen. Chemische Verbindungen, die die gleichen funktionellen Gruppen tragen, werden auf Grund ihrer oft ähnlichen Eigenschaften zu Stoffklassen zusammengefasst. Es gibt eine ganze Reihe solcher Stoffklassen, die in Abb. 18 dargestellt werden. Eine sei beispielhaft erwähnt.

TABLE 25.4 Common Functional Groups in Organic Compounds

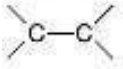
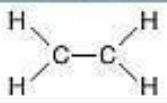
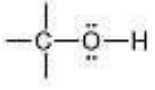
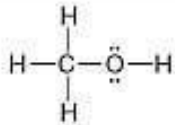
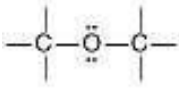
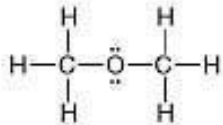
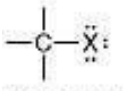
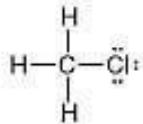
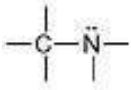
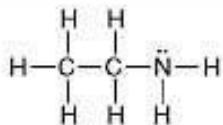
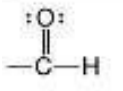
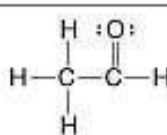
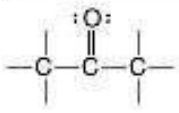
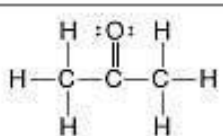
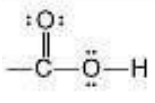
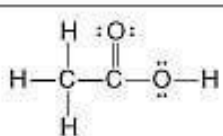
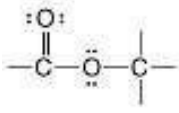
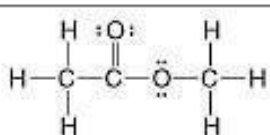
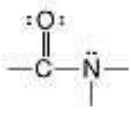
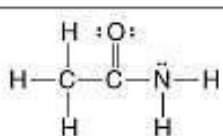
Functional Group	Type of Compound	Suffix or Prefix	Example	Systematic Name (common name)
	Alkene	-ene		Ethene (ethylene)
$\text{—C}\equiv\text{C—}$	Alkyne	-yne	$\text{H—C}\equiv\text{C—H}$	Ethyne (acetylene)
	Alcohol	-ol		Methanol (methyl alcohol)
	Ether	ether		Dimethyl ether
 (X = halogen)	Haloalkane	halo-		Chloromethane (methyl chloride)
	Amine	-amine		Ethylamine
	Aldehyde	-al		Ethanal (acetaldehyde)
	Ketone	-one		2-Propanone (acetone)
	Carboxylic acid	-oic acid		Ethanoic acid (acetic acid)
	Ester	-oate		Methyl ethanoate (methyl acetate)
	Amide	-amide		Ethanamide (acetamide)

Abbildung 18 funktionelle Gruppen

Jeder kennt Alkohol, doch es gibt verschiedene Alkoholarten. Derjenige, den wir gerne im Bier, Wein oder Wodka haben, ist Ethanol (Abb. 19). Ethanol besteht aus zwei Kohlenstoffatomen, 6 Wasserstoffatomen und einem Sauerstoffatom. Die sogenannte Summenformel, also wenn ich alle Atome zusammenzähle, wäre C_2H_6O . Doch tatsächlich wird am häufigsten die Summenformel C_2H_5OH verwendet; also wird ein Wasserstoffatom (H) gesondert hinter dem Sauerstoffatom gestellt. Warum? Das OH stellt nämlich die funktionelle Gruppe der Alkohole dar; an dieser OH-Gruppe erkennt man also die Alkohole, stellen sie in diese Stoffklasse mit den entsprechenden Eigenschaften. In der offiziellen Namensgebung der Chemie werden Alkohole mit der Endung „-ol“ genannt; wie bei Ethanol. Es gibt natürlich eine Reihe weiterer Alkohole; bekannt ist z. B. das Methanol, der einfachste Alkohol mit der Summenformel CH_3OH . Es hat also anders als das Ethanol nur einen Kohlenstoff. Aber dieser kleine Unterschied hat enorme Effekte, denn Methanol ist in homöopathischen Dosen tödlich. Würden also Wunderheiler in ihre Globuli Methanol reintun, hätten diese tatsächlich einen Effekt, wenn auch keinen positiven.

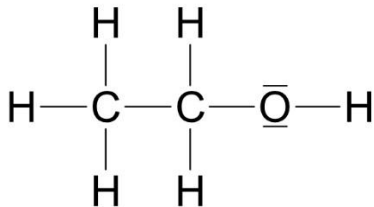
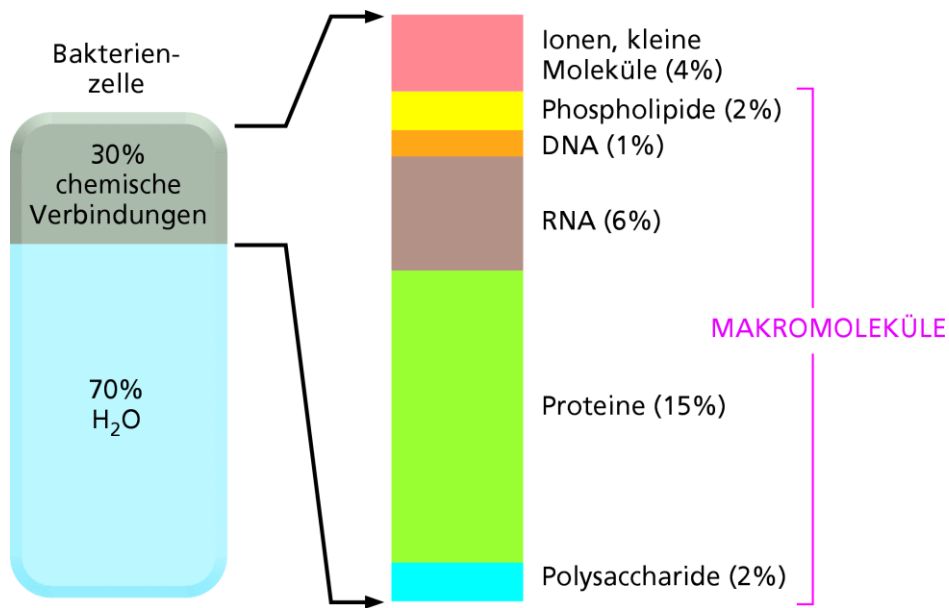


Abbildung 19 Ethanol

Die wichtigsten Biomoleküle die wir haben sind Proteine, Kohlenhydrate, Fette und Nukleinsäuren. Wir werden uns in diesem Kapitel mit den ersten drei beschäftigen, die Nukleinsäuren, wozu auch die DNS gehört, verdient eine gesonderte Folge.

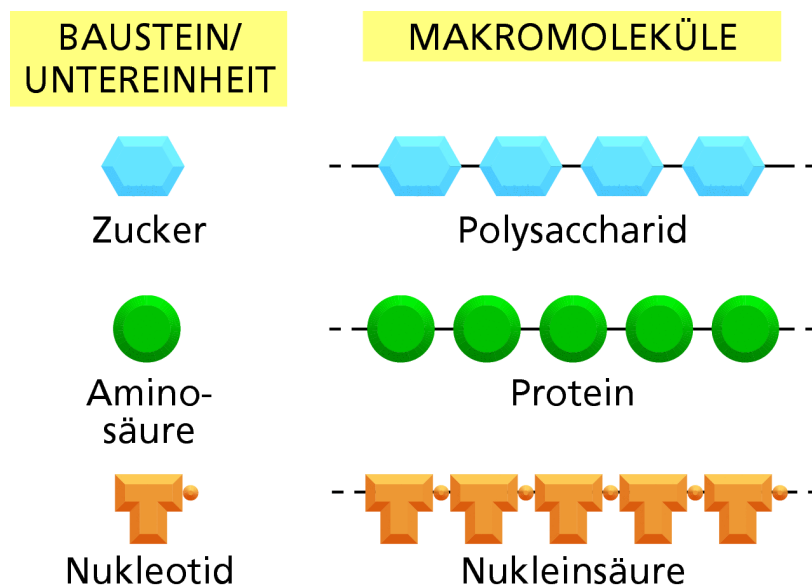
Eine Zelle, in Abb. 20 dargestellt am Beispiel einer Bakterienzelle, besteht zu 70% aus Wasser. Die restlichen 30% machen verschiedene Moleküle aus: davon u. a. 15% Proteine, jeweils 2% Fette und Kohlenhydrate, 7% Nukleinsäuren (DNS und RNA) und 4% besteht aus Ionen (z. B. Salzen) und anderen kleinen Molekülen.



© 2012 Wiley-VCH, Weinheim
 Alberts - Lehrbuch der Molekularen Zellbiologie
 ISBN: 978-3-527-32824-6 Fig. 02-026

Abbildung 20 Makromoleküle kommen in Zellen sehr häufig vor. Hier ist die ungefähre Zusammensetzung einer Bakterienzelle dargestellt. Die Zusammensetzung einer tierischen Zelle ist ähnlich.

Unsere Biomoleküle sind Makromoleküle, die aus einzelnen Bausteinen bzw. Untereinheiten aufgebaut sind. Man bezeichnet sie auch als Monomere. Monomere verbinden sich dann durch kovalente Bindungen zu Polymeren, also den Makromolekülen (Abb. 21). Wir werden nun die Monomere und Polymere unserer Biomoleküle kennenlernen.



© 2012 Wiley-VCH, Weinheim
 Alberts - Lehrbuch der Molekularen Zellbiologie
 ISBN: 978-3-527-32824-6 Fig. 02-027

Abbildung 21 Polysaccharide, Proteine und Nukleinsäuren werden aus monomeren Untereinheiten gebildet. Jedes Makromolekül ist ein Polymer aus kleinen Molekülen (genannt Monomere, Bausteine oder Untereinheiten), die durch kovalente Bindungen miteinander verbunden sind.

Aminosäuren

Proteine machen den größten Anteil unserer Makromoleküle aus und sie bestehen aus Aminosäuren. Aminosäuren sind also die Monomere, Proteine die Polymere.

Wie ist der Grundaufbau einer Aminosäure? Schauen wir uns dazu Abb. 22 an:

Amino Acid Structure

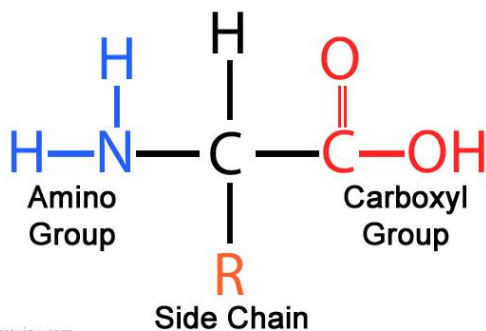


Abbildung 22 Grundaufbau der Aminosäure.

Zum einen haben wir ein Kohlenstoffatom, es ist das zentrale Kohlenstoffatom. An diesem befinden sich 4 Bindungsmöglichkeiten. Am oberen Ende ist eine Verbindung mit dem Wasserstoff. Auf der linken Seite findet sich eine Aminogruppe, die aus einem Stickstoffatom und zwei Wasserstoffatomen besteht. Rechts befindet sich die Carboxylgruppe. Eine Carboxylgruppe ist ein Kohlenstoffatom, dass eine Doppelbindung mit einem Sauerstoff eingeht und eine Verbindung mit einer Alkoholgruppe (Sauerstoff + Wasserstoff). Die Aminogruppe und die Carboxylgruppe sind die zwei funktionellen Gruppen der Aminosäure, die jede besitzt. Die vierte Bindung des zentralen Kohlenstoffatoms wird mit R abgekürzt. Sie steht für Rest bzw. Seitenkette. Diese ist bei verschiedenen Aminosäuren unterschiedlich aufgebaut. Die einfachste Aminosäure ist Glycin dessen Seitenkette nur aus einem Wasserstoffatom besteht.

Insgesamt kennt man über 250 Aminosäuren, für biologische Systeme sind aber nur 20 relevant, die die Proteine bilden.

Aminosäuren sind übrigens auch Zwitterionen (Abb. 23). Das Spannende ist nämlich, dass die Aminogruppe (NH₂) eine leicht basische, die Carboxylgruppe eine leicht saure Funktion hat. Das Vorhandensein beider solch gegensätzlicher Gruppen führt unweigerlich zu Wechselwirkungen. Dabei gibt nämlich die Carboxylgruppe ein Proton an die Aminogruppe ab. Dadurch wird die Carboxylgruppe negativ geladen (COO⁻) und die Aminogruppe positiv (NH₃⁺). Eine Aminosäure hat somit zwei Ladungen, die sich aber gegenseitig aufheben, wenn sich die Aminosäure im isoelektrischen Punkt befindet. Das ist ein exakt definierter pH-Wert einer wässrigen Lösung, bei dem sich bei Zwitterionen die positive und negative Ladung ausgleichen. Dieser pH-Wert ist von Aminosäure zu Aminosäure unterschiedlich, abhängig von der Rest-Seitenkette. Durch diese Ladungsverteilung wird die Aminosäure wasserlöslich.

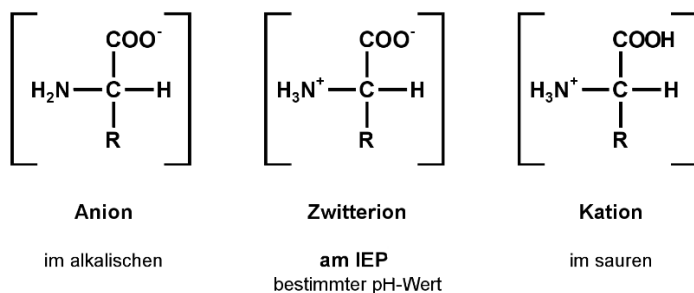


Abbildung 23 Zwitterionen der Aminosäuren

Innerhalb der organischen Chemie gibt es noch das Phänomen der Isomerie. Isomerie bedeutet, dass die Summenformel des Moleküls die gleiche ist, die Strukturformel jedoch eine andere. Das heißt im Grunde genommen nur, dass ein Molekül die gleiche Anzahl von Atomen haben kann, diese aber unterschiedlich angeordnet sein können. Das kann man z. B. beim Propanol, einem Alkohol mit drei Kohlenstoffatomen, ansehen (Abb. 24). Alkohole haben als funktionelle Gruppe die OH-Gruppe.

Beim Propanol kann aber die OH-Gruppe entweder am mittleren C-Atom oder am ersten C-Atom sein. Die Zahl der Atome ist die gleiche. Die Position der funktionellen Gruppe jedoch eine andere. Dadurch haben die Moleküle auch verschiedene chemische und physikalische Eigenschaften. Solche Isomere werden vor allem bei den Kohlenhydraten wichtig sein – schon als Vorwarnung.

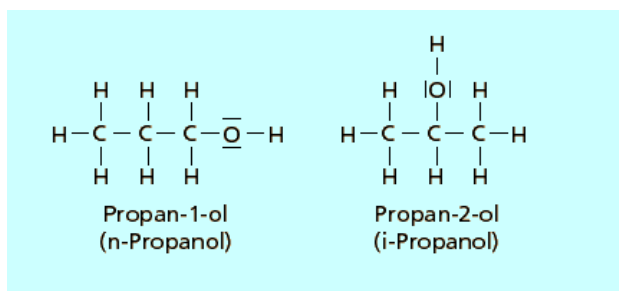


Abbildung 24 Isomere von Propanol

Eine weitere Form der Isomerie ist die Chiralität, auch Händigkeit genannt. Vielleicht hat jemand schon von linksdrehenden und rechtsdrehenden Aminosäuren etwas gehört. Das ist die besagte Chiralität. Das ist wie mit unseren Händen vergleichbar: dieselbe Anzahl an Fingern, Knochen, Muskeln etc. Und doch können wir unsere Hände in eine linke und eine rechte Hand unterscheiden. Vereinfacht ausgedrückt: links ist, wo der Daumen rechts ist. Dasselbe kann auch auf Moleküle zutreffen, doch da geht es nicht um Finger, sondern um die Lage funktionellen Gruppen. Bei den Aminosäuren kommt es darauf an, an welcher Seite die Aminogruppe und die Carboxylgruppe liegen (und sich im dreidimensionalen Raum das Molekül natürlich bewegt). Abb. 25 zeigt dies an einem Beispiel:

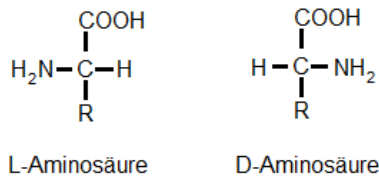


Abbildung 25 Chiralität der Aminosäuren

Befindet sich die Aminogruppe in dieser Darstellung (das ist die sog. Fischer-Projektion) links, haben wir die L-Form (von *levo*, links), steht sie rechts haben wir die D-Form (von *dextro*, rechts). Die Fischer-Projektion ist eine Methode, die Raumstruktur einer linearen, chiralen chemischen Verbindung eindeutig zweidimensional abzubilden.

Da es sich um gleiche Moleküle handelt, die sich lediglich im räumlichen Aufbau unterscheiden, haben sie auch identische chemische und physikalische Eigenschaften, sie unterscheiden sich aber in der Brechung und Polarisierung des Lichts. Aber so wie eine linke Hand nur in einen linken Handschuh passt, können sich beide Formen, L- und D-Aminosäuren nicht ineinander überführen. Lebewesen bevorzugen daher eine der Formen. Proteine sind bei Lebewesen fast ausschließlich aus L-Aminosäuren aufgebaut. Diese Übereinstimmung bei allen Lebewesen wird auch als ein wichtiger Hinweis für den gemeinsamen Ursprung allen Lebens gedeutet. Lediglich bei einigen Bakterien hat man in ihrer Zellwand einige Proteine mit rechtsdrehenden Aminosäuren gefunden, möglicherweise als Anpassung gegen bestimmte Antibiotika, die die Zellwände der Bakterien zerstören.

Wie schon erwähnt sind 20 verschiedene Aminosäuren am Aufbau unserer Proteine beteiligt, die sich in ihrer Seitenkette unterscheiden. Je nachdem aus welchen Atomen diese Seitenketten bestehen, kann man sie in verschiedene Untergruppen einteilen (Abb. 26):

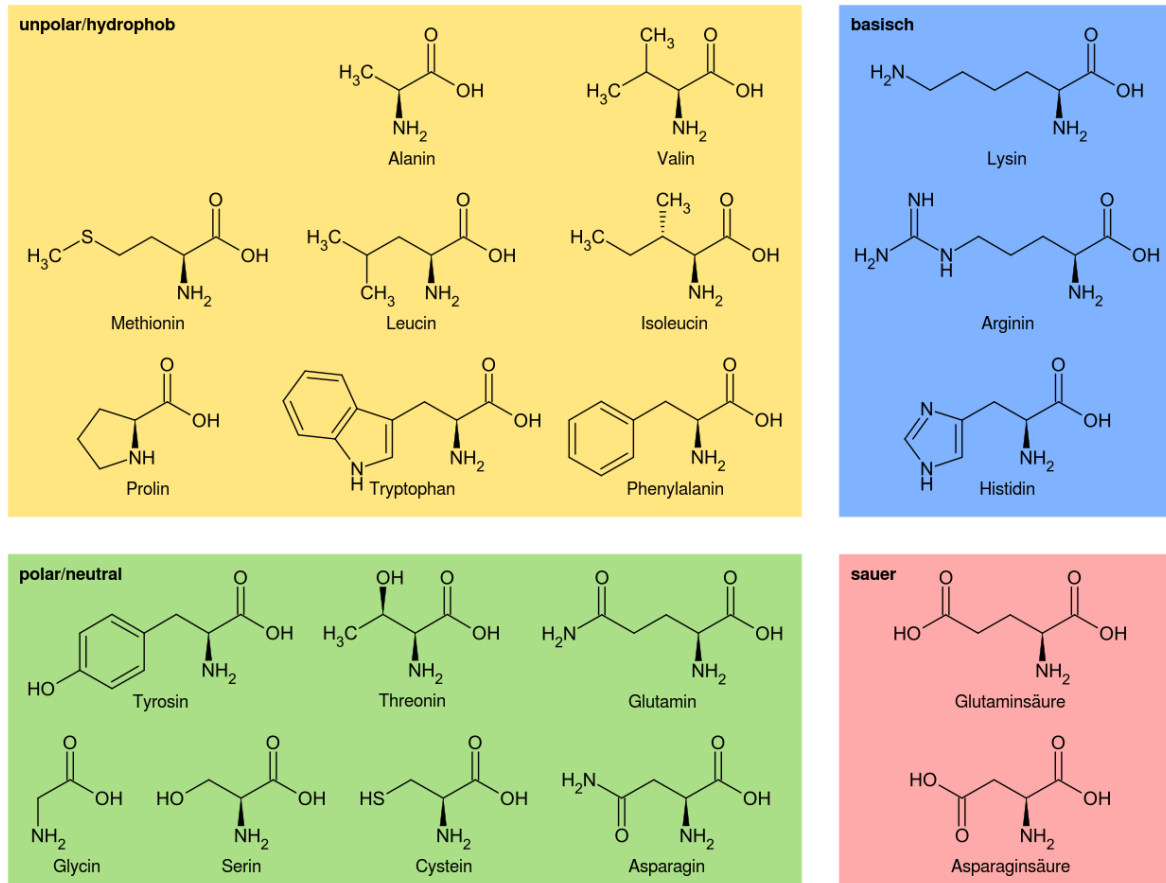


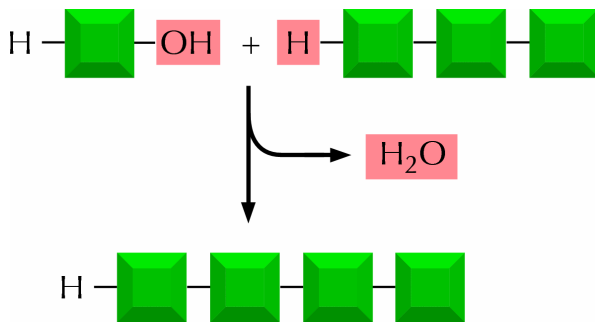
Abbildung 26 die 20 Aminosäuren

- unpolar und neutral, dazu gehören Aminosäuren wie Alanin oder Leucin
- polar und neutral, dazu gehören Aminosäuren wie Glyzin und Glutamin
- sauer, dazu gehören Aminosäuren wie die Asparaginsäure
- basisch, dazu gehören Aminosäuren wie Arginin oder Lysin

Wie die Einteilung vermuten lässt, haben diese Gruppen der Aminosäuren verschiedene chemische Eigenschaften, die sich dann auf die Eigenschaften der Proteine auswirken.

Wenn Aminosäuren eine Verbindung eingehen, spricht man von Peptidbindung. Verbinden sich zwei Aminosäuren, hat man ein Dipeptid, bei drei Aminosäuren ein Tripeptid ... und irgendwann zählt man nicht mehr so genau und man hat Polypeptide. Proteine sind Polypeptide. Wie läuft eine Peptid-Bindung ab? Die Carboxylgruppe von Aminosäure 1 geht eine Verbindung mit der Aminogruppe von Aminosäure 2 ein. D. h. beide Enden reagieren miteinander, dabei spaltet sich ein Wassermolekül ab und die Amino- und Carboxylgruppe verbinden sich miteinander, bilden eine Peptidbindung (Abb. 27). Hier kann man sich merken: wenn Aminosäuren, Kohlenhydrate oder andere Monomere sich zu Polymeren vereinigen, geschieht dies unter Wasserabspaltung. Man spricht auch von Kondensationsreaktion. An der Verbindung sind Enzyme beteiligt, die die Monomere miteinander verbinden und dieser Prozess kostet Energie. Die umgekehrte Reaktion – der Abbau des Polymers – erfolgt durch die Aufnahme von Wasser (Hydrolyse) und wird ebenfalls von Enzymen unter

Energieverbrauch durchgeführt. Das werden wir genauer im nächsten Beitrag behandeln.



© 2012 Wiley-VCH, Weinheim
Alberts - Lehrbuch der Molekularen Zellbiologie
ISBN: 978-3-527-32824-6 Fig. 02-028

Abbildung 27 Makromoleküle werden durch die Anlagerung von Untereinheiten an ein Ende gebildet. In einer Kondensationsreaktion wird pro Anlagerung eines Monomers an ein Ende der wachsenden Kette ein Molekül Wasser abgespalten. Die umgekehrte Reaktion – der Abbau des Polymers – erfolgt durch die Aufnahme von Wasser (Hydrolyse).

Proteine

Aminosäuren sind also die Grundbausteine der Proteine. Ein Protein jedoch nur als eine lange Kette von Aminosäuren zu definieren, wird diesem nicht gerecht.

Es gibt zwei Proteinklassen: Faserproteine und globuläre Proteine (Abb. 28). Faserproteine gibt es nur bei Tieren. Sie dienen oft als Strukturbildner, z. B. in Form von Bindegewebe oder Muskelfasern. Globuläre Proteine dienen für gewöhnlich nicht der Formgebung, sondern sind hauptsächlich für den Stoffwechsel und Transport zuständig. Membranproteine, Enzyme oder z. B. das Hämoglobin gehören zu den globulären Proteinen.

Globular and Fibrous Proteins

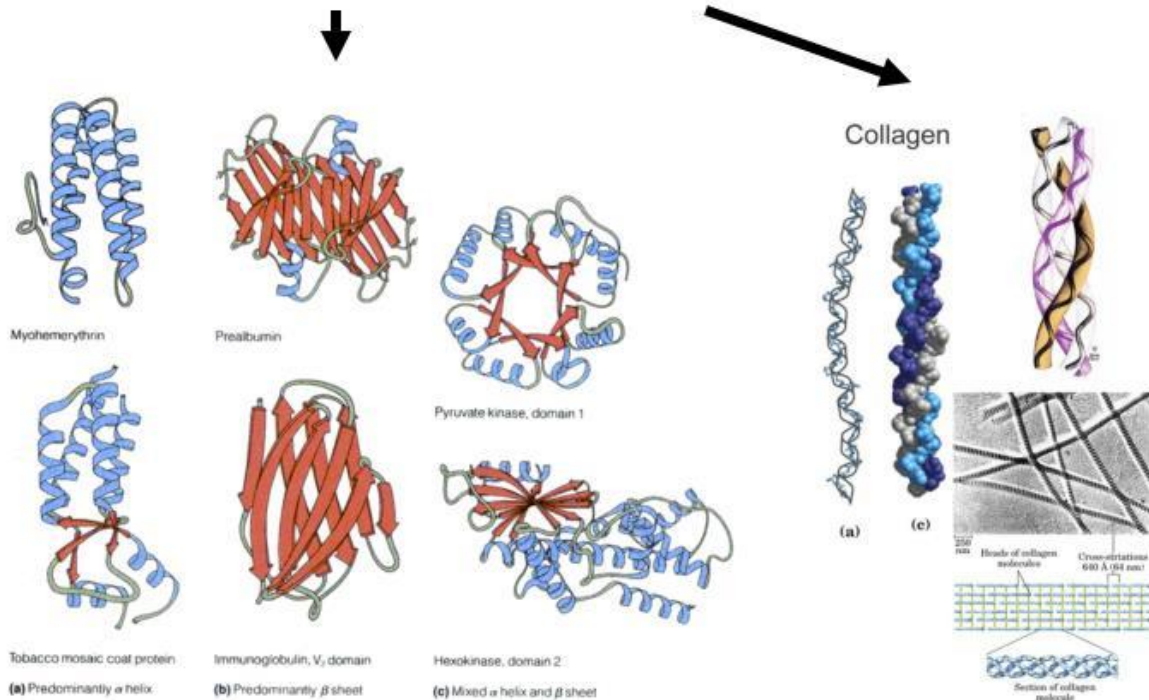


Abbildung 28 globuläre und Faserproteine

Proteine erfüllen damit eine wichtige Funktion bei Lebewesen. Faserproteine dienen der Struktur, Enzyme spalten Nährstoffe, Muskelproteine sorgen für die Bewegung, andere Proteine wie das Hämoglobin für den Transport anderer Moleküle, andere Proteine dienen der Regulation der Genaktivität, andere wie Antikörper dem Schutz und wieder andere für die Wahrnehmung, so z. B. das Rhodopsin, welches für den Sehprozess verantwortlich ist.

Die Funktion eines Proteins ist eng mit seiner Struktur verbunden. Nach der Struktur der Proteine lassen sich diese auf vier verschiedenen Ebenen betrachten: Primärstruktur, Sekundärstruktur, Tertiärstruktur und Quartärstruktur (Abb. 29).

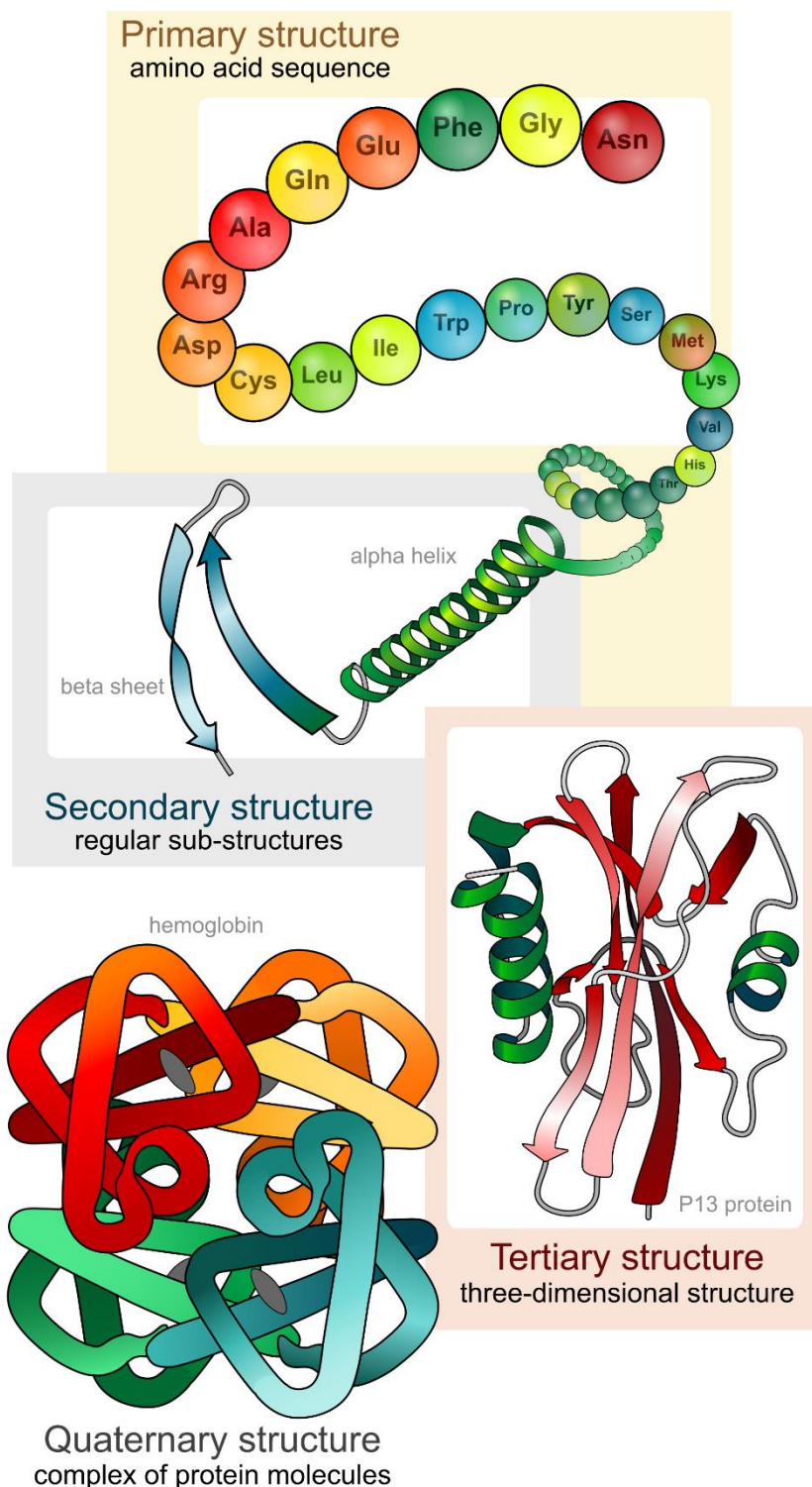


Abbildung 29 Die dreidimensionale Struktur der Proteine

Die Primärstruktur beschreibt die Abfolge der Aminosäuren in einem Protein.

Die Sekundärstruktur entsteht durch Wasserstoffbrückenbindungen zwischen benachbarten Proteinbereichen. Die bekanntesten Sekundärstrukturen sind die α -Helix und das β -Faltblatt.

Tertiärstrukturen entstehen durch Bindungen oder Wechselwirkungen in weiter entfernt liegenden Bereichen eines Proteins.

Eine Quartärstruktur besitzen nur Proteine, die aus mehreren Untereinheiten bestehen, wie beim Hämoglobin.

Die Sekundär-, Tertiär- und Quartärstrukturen lassen sich relativ einfach zerstören; z. B. durch Erhitzen. Jeder, der ein Ei gekocht oder ein Schnitzel gebraten hat, weiß das.

Kohlenhydrate

Die zweite Gruppe von Biomolekülen, mit denen wir uns näher befassen, sind die Kohlenhydrate. Zu den Kohlenhydraten gehören die verschiedenen Zucker, Stärke, Cellulose usw. Kohlenhydrate sind das Produkt der Photosynthese bei Pflanzen und einigen Bakterien, die aus CO_2 mithilfe von Lichtenergie Kohlenhydrate herstellen können. Kohlenhydrate sind die zentrale Energiequelle, auf der fast alle Stoffwechselprodukte beruhen. Kohlenhydrate sind eine Verbindung von Kohlenstoff, Wasserstoff und Sauerstoff.

Chemisch handelt es sich bei Kohlenhydraten um Oxidationsprodukte mehrwertiger Alkohole, also Alkohole, die mehr als eine OH-Gruppe als funktionelle Gruppe haben (unser Trinkalkohol Ethanol ist ein einwertiger Alkohol, weil er nur eine OH-Gruppe hat). Hierbei unterscheidet man Aldosen, z. B. Glucose) und Ketosen, z. B. Fructozucker (Abb. 30).

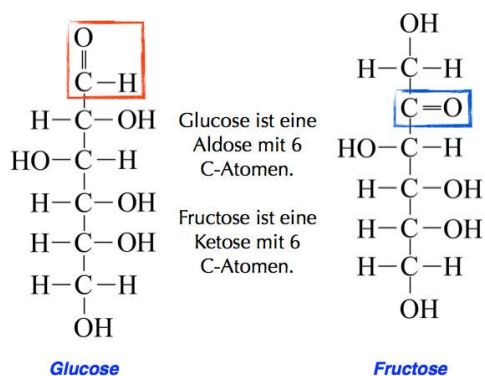


Abbildung 30 Aldosen und Ketosen beim Zucker

Glucose und Fructose sind beides Zucker mit 6 C-Atomen. Glucose als Aldose hat eine Aldehyd-Bindung als funktionelle Gruppe, bestehend aus einer Doppelbindung von Kohlenstoff und Sauerstoff und zusätzlich hat das Kohlenstoffatom eine weitere Bindung mit einem Wasserstoff. Fructose als Ketose hat eine Ketogruppe als funktionelle Gruppe. Das ist eine Doppelbindung des Kohlenstoffs mit einem Sauerstoff.

Auch bei Kohlenhydraten gibt es isomere Formen, je nachdem an welcher Stelle sich die OH-Gruppen befinden. Zusätzlich haben Kohlenhydrate auch eine Chiralität. Nehmen wir das Beispiel der C6-Aldose, dessen bekanntestes Isomer die Glucose ist (Abb. 31).

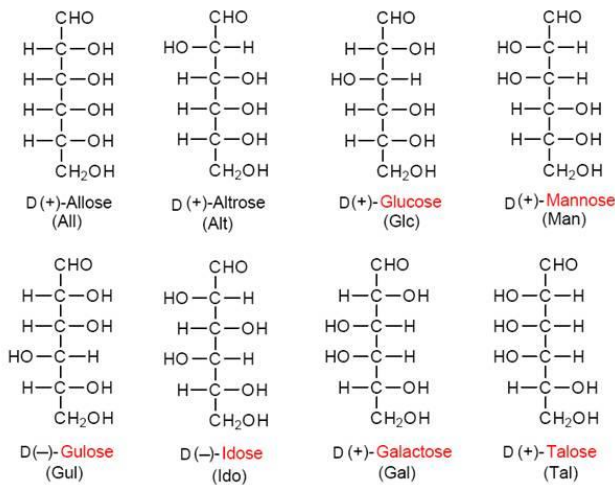


Abbildung 31 Isomere der D-C6-Aldosen

Man unterscheidet bei einer C6-Aldose 16 Varianten, je nachdem wo sich die OH-Gruppen befinden. Die 16 Isomere teilen sich auf in 8 Strukturisomere, die jeweils eine Chiralitätsform haben. Anders als bei den Aminosäuren dominieren bei Lebewesen die D-Formen. Die bekanntesten C-6 Aldosen sind die D-Glucose, D-Galactose und D-Mannose. Unter den Ketosen ist die D-Fructose unser wichtigstes.

Neben C-6-Kohlenhydraten gibt es auch einige wichtige, die nur 5 Kohlenstoffatome haben. Hierzu zählen die Ribose und Desoxyribose, die das Grundgerüst von RNA und DNA ausmachen. Zucker mit 5 C-Atomen nennt man Pentose, Zucker mit 6 C-Atomen nennt man Hexosen.

Pentosen und Hexosen neigen dazu Ringformen zu bilden (Abb. 32). Bei solch einer Ringbildung kann die OH-Gruppe des ersten C-Atoms in zwei verschiedene Richtungen zeigen. Zeigt bei der Ringform die OH-Gruppe des ersten C-Atoms nach unten, haben wir eine alpha-Form, zeigt die OH-Gruppe nach oben, haben wir die beta-Form.

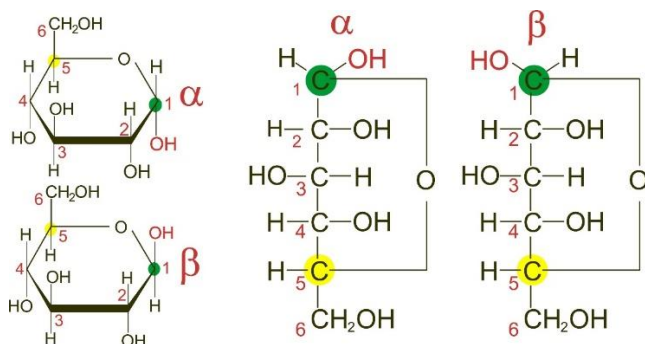


Abbildung 32 Ringbildung bei der D-Glucose

Die Monomere der Kohlenhydrate bezeichnet man als Monosaccharide (Einfachzucker). Verbinden sich zwei Monosaccharide (unter Wasserabspaltung, wie bei den Aminosäuren) werden sie zu Disacchariden. Kohlenhydrate mit bis zu 10 Monomeren bezeichnet man als Oligosaccharide, mit mehr als 10 Monomeren als Polysaccharide. Verbinden sich Zuckermoleküle miteinander, spricht man von glykosidischer Bindung. Eine glykosidische Bindung wird immer mit 2 Ziffern angegeben, welche die beiden C-Atome angeben, zwischen denen, die Bindung besteht. Beispiel Maltose ist ein Disaccharid mit einer 1,4 glykosidischen Bindung zwischen zwei α -D-Glucose (Abb. 33).

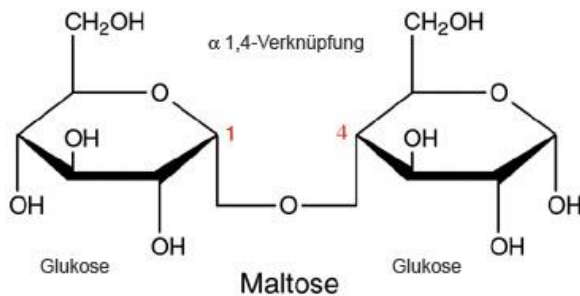


Abbildung 33 1,4 glykosidischen Bindung zwischen zwei α -D-Glucose führen zum Disaccharid Maltose

Die Monosaccharide, Disaccharide und Oligosaccharide sind in der Regel wasserlöslich, haben einen süßen Geschmack und werden im engeren Sinne als Zucker bezeichnet. Die Polysaccharide sind hingegen oftmals schlecht oder gar nicht in Wasser löslich und geschmacksneutral.

Wir haben ja einige Einfachzucker kennengelernt.

Die wichtigsten Zweifachzucker sind:

Saccharose (Rübenzucker, Glucose & Fructose)

Lactose (Milchzucker, Glucose + Galaktose)

Maltose (Malzzucker, Glucose + Glucose)

Die wichtigsten Polysaccharide sind Stärke, Glykogen (in unserer Leber gespeichert), Cellulose und Chitin.

Stärke und Cellulose bestehen beide aus D-Glucose, wobei Stärke aus α -1,4- und α -1,6 glykosidischen Verbindungen besteht und Cellulose aus β -1,4-Bindungen.

Lipide

Lipide sind chemisch heterogene Substanzen, die sich schlecht in Wasser (Hydrophobie), gut dagegen in unpolaren Lösungsmitteln (Lipophilie) lösen. In lebenden Organismen werden Lipide hauptsächlich als Strukturkomponente in Zellmembranen, als Energiespeicher oder als Signalmoleküle gebraucht. Die meisten biologischen Lipide sind amphiphil, besitzen also einen lipophilen Kohlenwasserstoff-Rest und eine polare hydrophile Kopfgruppe, deshalb bilden sie in polaren Lösungsmitteln wie Wasser Micellen oder Membranen. Oft wird der Begriff Fett als Synonym für Lipide gebraucht, jedoch stellen die Fette nur eine Untergruppe der Lipide dar (nämlich die Gruppe der Triglyceride).

Die Lipide können in sieben Gruppen eingeteilt werden: Fettsäuren, Triglyceride (Neutralfette), Wachse, Phospholipide, Sphingolipide, Glycolipide und Isoprenoide (Steroide, Carotinoide etc. Abb. 34).

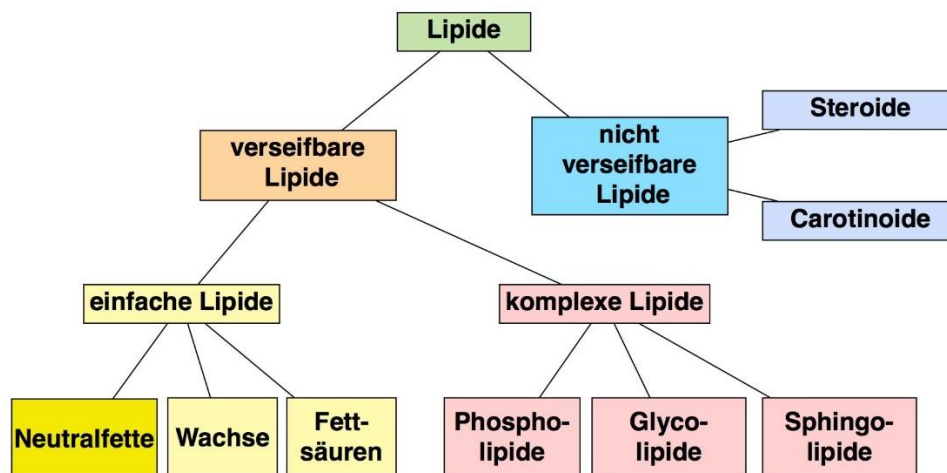


Abbildung 34 Einteilung der Lipide

Neutralfette bzw. Triglyceride machen mit >90 % den Hauptanteil der Nahrungslipide aus. Sie sind ein wichtiger Energielieferant (1 g Fett enthält 39 kJ Energie, 1 g Zucker nur 17 kJ). Außerdem bilden sie den wichtigsten Energiespeicher des Körpers (Zucker, d. h. Glucose, wird dagegen in viel geringerer Menge als Glycogen in der Leber gespeichert), sie sind ein guter Kälteschutz in der Haut und schützen diese auch vor Verletzungen. Alle wichtigen Organe werden durch einen Fettmantel geschützt. Triglyceride bestehen aus dem dreiwertigen Alkohol Glycerin und drei Fettsäuren (Abb. 35):

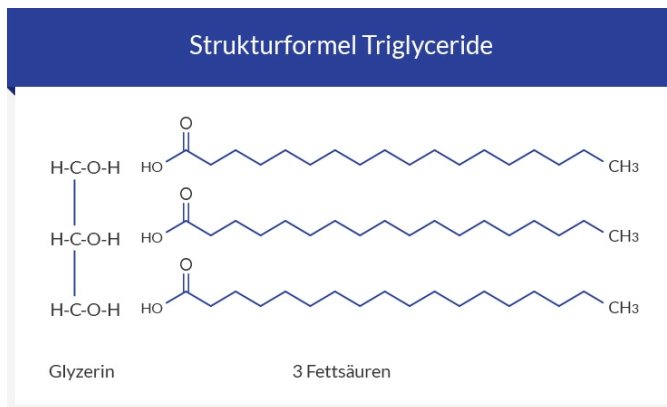


Abbildung 35 Struktur eines Triglycerids

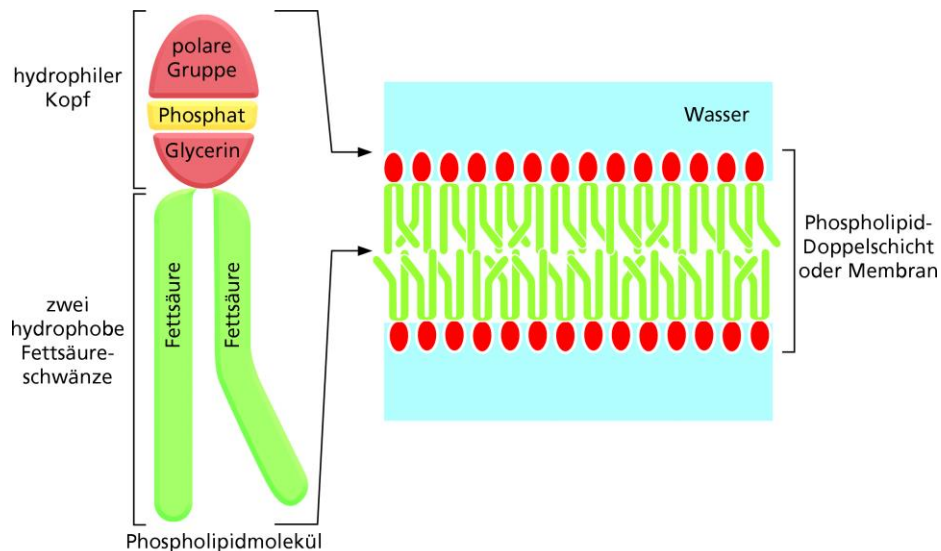
Fettsäuren sind meist unverzweigte Monocarbonsäuren, die aus einer Kette aus Kohlenstoffatomen bestehen, an deren einem Ende sich eine Carboxylgruppe befindet.

Unterschieden wird zwischen gesättigten Fettsäuren, in denen keine Doppelbindungen vorkommen, und ungesättigten Fettsäuren, die eine oder mehrere Doppelbindungen besitzen. Die einfachste gesättigte Fettsäure ist die Buttersäure und enthält nur vier Kohlenstoffatome. Ein wichtiger Vertreter der ungesättigten Fettsäuren ist z. B. die Ölsäure (einfach ungesättigt) und Arachidonsäure (vierfach ungesättigt). Je mehr Doppelbindungen Fettsäuren enthalten, desto niedriger liegt ihr Schmelzpunkt, sind also bei Raumtemperatur eher flüssig, z. B. bei Ölen. Ungesättigte Fettsäuren können vom tierischen Organismus nur unter Einschränkung synthetisiert werden. Man bezeichnet daher all jene Fettsäuren, die mit der Nahrung aufgenommen werden müssen als 'essenzielle Fettsäuren'

Wachse unterscheiden sich von den Fetten dadurch, dass anstelle des Glycerins andere Alkohole als Grundstruktur eintreten (z. B. Myricylalkohol). Wachse sind ebenso wie Fette neutrale Verbindungen, die unpolare langkettige Kohlenwasserstoffreste enthalten. Daher sind sie nicht in Wasser löslich.

Neben tierischen Wachsen (z. B. Bienenwachs, Walrat) sind auch Pflanzenwachse bekannt, z. B. das Carnauba-Wachs, das für Autopolituren Verwendung findet.

Membranbildende Lipide sind Lipide, die einen hydrophilen und einen hydrophoben Teil besitzen – also amphiphil sind (Abb. 36). Dies erlaubt es ihnen, in polaren Lösungsmitteln wie Wasser je nach Beschaffenheit entweder Mizellen (kugelförmige Aggregate aus amphiphilen Molekülen, die sich spontan zusammenlagern) oder Doppellipidschichten bilden – wobei immer der hydrophile Teil mit dem polaren Lösungsmittel interagiert. Aus diesen Doppellipidschichten sind alle Biomembranen aufgebaut, welche den Inhalt einer Zelle gegen die Umgebung abgrenzt und somit membranbildende Lipide zu einer der Grundvoraussetzungen allen Lebens macht. Phospholipide haben eine Phosphatgruppe und bilden den Hauptbestandteil von Biomembranen. In diesen Substanzen ist Glycerin mit zwei Fettsäuren und Phosphorsäure verestert. Der Phosphorsäurerest besitzt eine negative elektrische Ladung und ist mit einem Alkoholrest (z. B. Ethanolamin, Cholin, Colamin, Serin, Inosit oder Glycerin) verbunden.



© 2012 Wiley-VCH, Weinheim
 Alberts - Lehrbuch der Molekularen Zellbiologie
 ISBN: 978-3-527-32824-6 Fig. 02-020

Abbildung 36: Phospholipide lagern sich zusammen und bilden so die Zellmembranen. Phospholipide bestehen aus zwei hydrophoben Fettsäureschwänzen, die mit einem hydrophilen Kopf verbunden sind. In wässriger Umgebung lagern sich die hydrophoben Schwänze der Fettsäurelipide zusammen, um Wasser auszuschließen. Dadurch entsteht eine Doppelschicht, bei der die hydrophilen Köpfe jedes Phospholipids zum Wasser gerichtet sind.

Sphingolipide sind komplex aufgebaute Phospholipide, die gehäuft im Gehirn und im Nervengewebe anzutreffen sind. Im wichtigsten Vertreter, dem Sphingomyelin, ist anstelle des Glycerinrests der Aminoalkohol Sphingosin enthalten.

Als Isoprenoide (auch Terpenoide) werden Verbindungen bezeichnet, die auf Isopreneinheiten aufbauen. Zu diesen Lipiden zählende Verbindungen sind die Steroide und die Carotinoide. Der bekannteste Vertreter der Steroide ist das zu den Sterinen zählende Cholesterin. Es ist unter anderem auch ein essentieller Bestandteil aller Zellmembranen mit Ausnahme der Innenmembran der Mitochondrien, und kann somit im erweiterten Sinne auch zu den Membranlipiden gezählt werden. Auch die Sexualhormone sind Steroide. Carotinoide sind Polymerisationsprodukte von Isopren, die ausschließlich in Pflanzen hergestellt werden und dort als gelb bis rötliche Farbstoffe fungieren.

Kapitel 4: chemische Reaktion, Katalyse, Enzyme und ATP

Im vorherigen Beitrag haben wir uns die verschiedenen Biomoleküle angesehen: Proteine, Kohlenhydrate und Lipide. Sie bilden den Grundaufbau unserer Zellen und liefern uns die Energie für unsere Lebenserhaltung. Sie sind also essentieller Bestandteil des Stoffwechsels.

Als Stoffwechsel oder Metabolismus bezeichnet man die gesamten chemischen und physikalischen Vorgänge der Umwandlung chemischer Stoffe bzw. Substrate in Zwischenprodukte und Endprodukte im Organismus von Lebewesen.

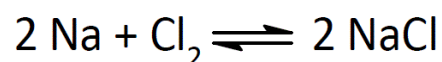
Chemische Reaktion, Oxidation & Reduktion

Chemisch betrachtet sind Stoffwechsel chemische Reaktionen. Eine chemische Reaktion ist ein Vorgang, bei dem eine oder meist mehrere chemische Verbindungen in andere umgewandelt werden und Energie freigesetzt oder aufgenommen wird.

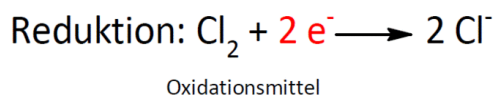
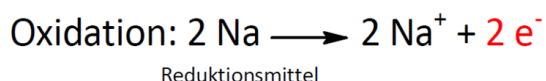
Zwei wichtige Begriffe hierbei sind Oxidation und Reduktion. Beide gehen immer mit einer Übertragung von Elektronen einher. Unter einer Reduktion versteht man die Aufnahme von Elektronen und eine Erniedrigung der Oxidationszahl.

Eine Oxidation ist eine Abgabe von Elektronen und eine Erhöhung der Oxidationszahl. Oxidation und Reduktion sind immer aneinandergekoppelt. Eine Oxidation kann nicht ablaufen, wenn die abgegebenen Elektronen nicht sofort in einer Reduktion aufgenommen werden können. Man kann zwar die beiden Prozesse räumlich voneinander trennen, jedoch müssen sie immer in Kontakt zueinanderstehen. Entsprechend spricht man von einer Redoxreaktion. Schauen wir uns das an einem einfachen Beispiel an.

Chemische Reaktionen werden in einer Reaktionsgleichung dargestellt. Links stehen die Ausgangsstoffe, die Edukte, rechts das Ergebnis, die Produkte. Bei einer Reaktionsgleichung müssen links und rechts die gleiche Anzahl an Atomen stehen. Wir haben links als Edukte zwei Natriumatome und eine Chlor-Verbindung, die aus zwei Chloratomen zusammengesetzt sind. Natrium und Chlor reagieren zu NaCl. Da auf der Seite der Edukte je zwei Natrium und zwei Chloratome zu finden sind, entstehen bei den Reaktionsprodukten 2 NaCl. Die chemische Reaktion wird durch einen Reaktionspfeil dargestellt. Wir sehen hier aber, dass es zwei Pfeile gibt: einer zeigt in Richtung Reaktionsprodukt, der andere in Richtung der Edukte. Es handelt sich hierbei also um eine Redoxreaktion, die Reaktion kann also in beide Richtungen ablaufen.



Schauen wir uns daher an, wie in dieser Reaktion die Oxidation und wie die Reduktion verläuft:



Wir wissen aus dem Beitrag über die Grundlagen der Chemie, dass Natrium im Periodensystem in der ersten Hauptgruppe steht. Schaut euch dieses Video nochmal an, um den Rest zu verstehen. Es hat also ein Elektron auf der äußersten Schale und um den Edelgaszustand zu erreichen, gibt es gerne sein äußerstes Elektron ab. D. h. hier findet eine Oxidation statt: Natrium gibt sein Elektron ab. Die Natriumatome liefern die Elektronen und werden somit oxidiert. Aus den Natriumatomen entstehen zwei Natriumionen und zwei Elektronen. Dabei wird das Natrium zum Reduktionsmittel. Reduktionsmittel sind Verbindungen, die durch Abgabe von Elektronen eine Reduktion einer anderen Verbindung bewirken. Dabei werden sie selbst oxidiert.

Chlor steht in der VII. Hauptgruppe im Periodensystem, hat also sieben Elektronen auf der Außenbahn. Dem Chlor fehlt nur noch ein Elektron zum Edelgaszustand. Dieses nimmt entsprechend Elektronen auf, wird also reduziert. Damit fungiert Chlor als Oxidationsmittel. Oxidationsmittel sind die Verbindungen, die durch Aufnahme von Elektronen eine Oxidation einer anderen Verbindung bewirken. Dabei werden sie selbst reduziert.

Fassen wir also zusammen: Natrium gibt Elektronen ab, wird also oxidiert und zum Reduktionsmittel (weil es für die Reduktion des Chlors sorgt). Chlor nimmt Elektronen auf, wird reduziert und agiert als Oxidationsmittel, weil es für die Oxidation des Natriums sorgt.

Oxidation = Abgabe von Elektronen

Reduktion = Aufnahme von Elektronen

Oxidationsmittel = Elektronenakzeptor

Reduktionsmittel = Elektronendonator

In Redoxprozessen müssen immer alle abgegebenen Elektronen wieder aufgenommen werden. Es gibt keine frei diffundierenden Elektronen im Reaktionsmedium.

Exotherm, endotherm und Aktivierungsenergie

Soweit zu den Redoxreaktionen. Eine chemische Reaktion zwischen Natrium und Chlor läuft relativ schnell ab und quasi von selbst. Bei den Biomolekülen ist es etwas komplizierter.

Jede chemische Reaktion ist nicht nur mit einer Stoffumwandlung verbunden, sondern auch mit einer Energieumwandlung. Diese Energieumwandlung "ermöglicht" es, unter energetischen Gesichtspunkten eine chemische Reaktion zu unterteilen. Je nachdem, ob Energie (in Form von Wärme) an die Umgebung abgegeben oder aufgenommen

wird, teilt man eine chemische Reaktion in eine exotherme oder endotherme Reaktion ein. Diese Einteilung in exotherm und endotherm hat aber keinen Einfluss auf die Aktivierungsenergie. In der Regel muss Aktivierungsenergie hinzugeführt werden, um die Reaktion zu aktivieren (so dass die Stoffe in einer messbaren Geschwindigkeit miteinander reagieren). Die Aktivierungsenergie ist also eine energetische Barriere, die bei einer chemischen Reaktion von den Reaktionspartnern überwunden werden muss. Wird diese Barriere überwunden, kommt es erst dann zur chemischen Reaktion. Allgemein gilt: Je niedriger die Aktivierungsenergie, desto schneller verläuft die Reaktion.

Sehen wir uns nun Beispiele für exotherme und endotherme Reaktionen an.

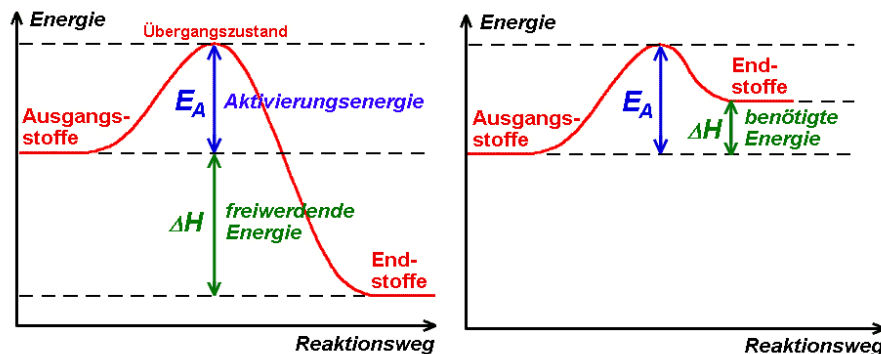
Exotherme Reaktion: Bei einer exothermen Reaktion wird Wärmeenergie an die Umgebung abgegeben. Eine typische exotherme Reaktion ist eine Verbrennungsreaktion. Der Ausgangsstoff, z. B. das Holz, das verbrannt werden soll, hat ein gewisses Energielevel. Aber das Holz alleine brennt nicht, es braucht eine Aktivierungsenergie. Das ist zum Beispiel die Energie, die du zuführst, wenn du das Holz anzündest. Dadurch fügst du dem System die Aktivierungsenergie zu und die Reaktion läuft ab: Das Holz verbrennt. Hierbei wird Energie freigesetzt in Form von Wärme, die Luft wird nämlich wärmer. Als Produkt der chemischen Reaktion entsteht die Holzkohle, das verbrannte Holz bzw. die Asche, die ein niedrigeres Energieniveau hat, als das Edukt Holz.

Endotherme Reaktion: Bei einer endothermen Reaktion wird Wärmeenergie aus der Umgebung aufgenommen. Daher laufen endotherme Reaktionen in der Regel nur solange ab, wie aus der Umgebung Energie zugeführt wird bzw. noch Ausgangsstoffe vorhanden sind. Stoppt bei einer endothermen Reaktion die Energiezufuhr aus der Umgebung, kommt in der Regel die Reaktion zum Stillstand. Ein Beispiel kennt ist die Braustablette. Beim Auflösen des Brausepulvers wird die Lösungsenergie für die Reaktion dem Wasser entzogen. Das Wasser kühlt sich bei diesem Vorgang etwas ab. Folglich wird Energie in Form von Wärme aus der Umgebung aufgenommen und die Reaktion ist damit endotherm.

Abb. 37 zeigt die Energiekurve von endothermen und exothermen Reaktionen.

Exotherme und endotherme Reaktionen

Energiediagramme unter Berücksichtigung der Aktivierungsenergie



exotherme Reaktion

Es wird Energie frei

endotherme Reaktion

Es wird Energie benötigt

www.seilnacht.tuttlingen.com

Abbildung 37 Energieniveau von endothermen und exothermen Reaktionen.

Enzyme

Dasselbe Prinzip läuft auch in unseren Zellen ab. Merkt euch aber eines alle chemischen Reaktionen brauchen eine Aktivierungsenergie. Aber hier stehen wir vor einem gewaltigen Problem: Gerade unsere Biomoleküle: ob wir sie nun abspalten oder aufbauen brauchen eine sehr hohe Aktivierungsenergie. Sich selbst überlassen bei Raumtemperatur oder gar Körpertemperatur reagieren unsere Proteine, Fette und Zucker überhaupt nicht, weil es eine hohe Aktivierungsenergie braucht. Man braucht also etwas, was die chemischen Reaktionen schneller ablaufen lassen würde. Man spricht hierbei von Katalysatoren. Das sind Stoffe, die eine chemische Reaktion beschleunigen können, ohne aber selbst Teil der chemischen Reaktion zu werden, also selbst umgewandelt zu werden. Katalysatoren setzen die Aktivierungsenergie herab und ermöglichen die chemische Reaktion. Katalysatoren kennt ihr vielleicht am besten aus Autos, der gesundheits- und umweltschädliche Abgase in ungiftige Stoffe umwandelt. Doch die meisten unserer Stoffwechselreaktionen brauchen auch Katalysatoren. Unsere Katalysatoren, zumindest die wichtigsten, sind die Enzyme.

Enzyme sind Proteine, die als Biokatalysator biochemische Reaktionen im Organismus steuern und beschleunigen, ohne dabei selbst verändert zu werden. Sie sind in allen Körperzellen enthalten und sind unerlässlich für alle Körperfunktionen. So steuern Enzyme nicht nur die Verdauung, sondern den gesamten Stoffwechsel und sind damit ein wesentlicher Faktor für die Gesundheit.

Enzyme binden spezifisch ein oder einige wenige Substrate, indem sie mit ihnen einen reversiblen Enzym-Substrat-Komplex bilden. Als Enzym-Substrat-Komplex bezeichnet man den Übergangszustand aus Enzym und gebundenem Substrat, der bei einer enzymkatalysierten Reaktion im ersten Schritt gebildet wird (Abb. 38).

Schlüssel-Schloss-Prinzip

- Enzyme können sich der Form ihres Substrates anpassen ("induced-fit").

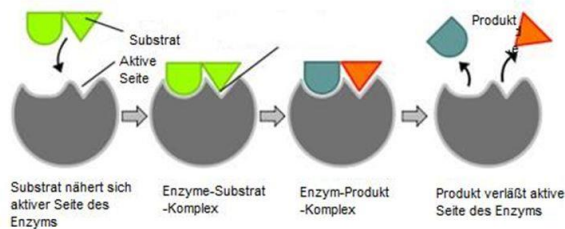


Abbildung 38 Enzym-Substrat-Komplex

Die spezielle Hohlstruktur im Enzym bewirkt, dass das aktive Zentrum mit einem passenden Substrat in Kontakt treten kann. Enzyme können dadurch nur ein bestimmtes oder einige wenige Substrate umsetzen, was Substratspezifität genannt wird. Des Weiteren ist ein Enzym nur in der Lage, eine bestimmte Reaktion zu katalysieren, was man Reaktionsspezifität nennt. Ein Substrat kann dagegen von unterschiedlichen Enzymen umgesetzt werden. Es funktioniert also nach dem Schlüssel-Schloss-Prinzip.

Bei der Reaktion ist die Substratbindungsstelle, das aktive Zentrum, von Bedeutung, also der Ort, an dem das Substrat gebunden wird. Oft können Enzyme ihre Wirkung erst dann entfalten, wenn sie von einem sogenannten Cofaktor oder mehreren Cofaktoren aktiviert werden. Das können Metallionen (wie Eisen-, Kupfer- oder Zink-Ionen) oder aber organische Moleküle (wie Vitamine) sein.

Enzyme lassen sich in sechs Hauptgruppen einteilen, je nachdem, welche Art von chemischer Reaktion sie katalysieren. Diese Enzymklassen (und einige ihrer Untergruppen) sind:

- 1. Oxidoreduktasen:** Sie katalysieren Reaktionen, bei denen Elektronen übertragen werden (Redoxreaktionen); z.B. Dehydrogenasen, Oxidasen, Reduktasen, Katalasen
- 2. Transferasen:** Sie katalysieren Reaktionen, bei denen ganze funktionelle Gruppen (wie Phosphat-Gruppe) von einem Molekül auf ein anderes übertragen werden; z.B. Transaminasen, Kinasen, DNA-Polymerasen
- 3. Hydrolasen:** Sie katalysieren Reaktionen, bei denen eine chemische Bindung entweder unter Wasseraustritt entsteht oder unter Wasseranlagerung gespalten wird; z.B. Peptidasen, Phosphatasen, Proteasen
- 4. Lyasen:** Sie katalysieren Reaktionen, bei denen chemische Bindungen ohne Energieverbrauch gespalten oder gebildet werden; z.B. Aldolase
- 5. Isomerasen:** Sie sorgen dafür, dass Bindungsverhältnisse innerhalb eines Moleküls neu geordnet werden; z.B. Racemasen, Topoisomerasen
- 6. Ligasen (Synthetasen):** Sie katalysieren Reaktionen, bei denen zwei Moleküle unter Energieverbrauch miteinander verbunden werden; z.B. Carboxylasen

Enzymkinetik

Die Enzymkinetik beschäftigt sich mit dem zeitlichen Verlauf enzymatischer Reaktionen. Eine zentrale Größe hierbei ist die Reaktionsgeschwindigkeit. Sie ist ein Maß für die Änderung der Substratkonzentration mit der Zeit, also für die Stoffmenge Substrat, die in einem bestimmten Reaktionsvolumen pro Zeiteinheit umgesetzt wird (Abb. 39).

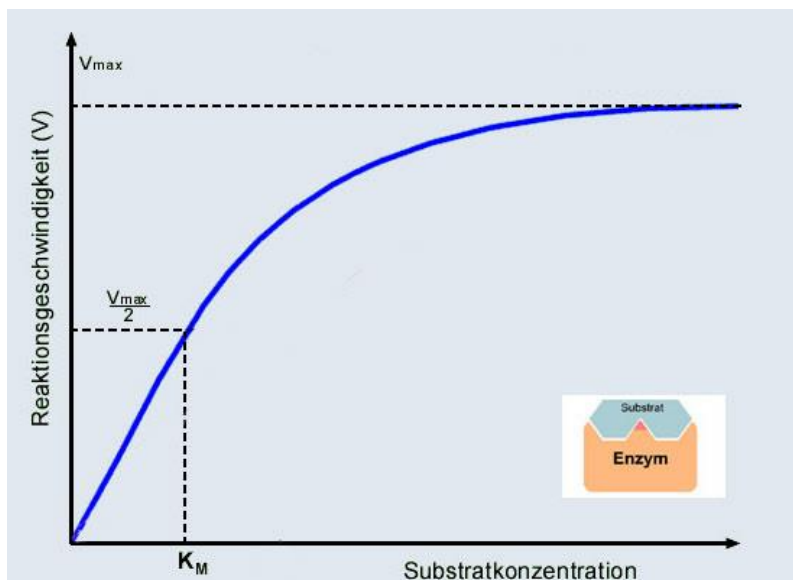


Abbildung 39 Reaktionsgeschwindigkeit von Enzymen

Neben den Reaktionsbedingungen wie Temperatur, Salzkonzentration und pH-Wert der Lösung, hängt sie von den Konzentrationen des Enzyms, der Substrate und Produkte sowie von Effektoren (Aktivatoren oder Inhibitoren) ab (Abb. 40).

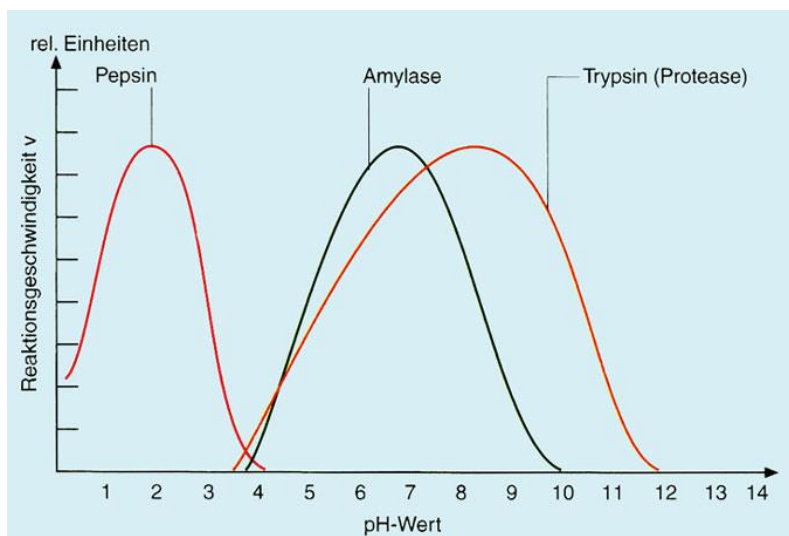


Abbildung 40: Enzymaktivität in Abhängigkeit vom pH-Wert

Die gemessene Enzymaktivität ist proportional zur Reaktionsgeschwindigkeit und damit stark von den Reaktionsbedingungen abhängig. Sie steigt mit der Temperatur

entsprechend der RGT-Regel an: eine Erhöhung der Temperatur um ca. 5-10 °C führt zu einer Verdoppelung der Reaktionsgeschwindigkeit und damit auch der Aktivität. Dies gilt jedoch nur für einen begrenzten Temperaturbereich. Bei Überschreiten einer optimalen Temperatur kommt es zu einem steilen Abfallen der Aktivität durch Denaturierung des Enzyms. Änderungen im pH-Wert der Lösung haben oft dramatische Effekte auf die Enzymaktivität, da dieser die Ladung einzelner für die Katalyse wichtiger Aminosäuren im Enzym beeinflussen kann (Abb. 40). Jenseits des pH-Optimums vermindert sich die Enzymaktivität und kommt irgendwann zum Erliegen.

Ein Modell zur kinetischen Beschreibung einfacher Enzymreaktionen ist die Michaelis-Menten-Theorie (Abb. 39 & 41).

Michaelis-Menten-Gleichung

$$V_o = \frac{V_{\max} [S]}{K_m + [S]}$$

V_o	Anfangsgeschwindigkeit
$V_{\max}$	Maximalgeschwindigkeit
$[S]$	Substratkonzentration
K_m	Michaelis-Menten-Konstante

Abbildung 41 Michaelis-Menten-Gleichung

Sie liefert einen Zusammenhang zwischen der Reaktionsgeschwindigkeit v einer Enzymreaktion sowie der Enzym- und Substratkonzentration $[E_0]$ und $[S]$. Grundlage ist die Annahme, dass ein Enzym mit einem Substratmolekül einen Enzym-Substrat-Komplex bildet und dieser entweder in Enzym und Produkt oder in seine Ausgangsbestandteile zerfällt. Was schneller passiert hängt von den jeweiligen Geschwindigkeitskonstanten k ab.

Das Modell besagt, dass mit steigender Substratkonzentration auch die Reaktionsgeschwindigkeit steigt. Das geschieht anfangs linear und flacht dann ab, bis eine weitere Steigerung der Substratkonzentration keinen Einfluss mehr auf die Geschwindigkeit des Enzyms hat, da dieses bereits mit Maximalgeschwindigkeit $V_{\max}$ arbeitet.

Enzyme können auch gehemmt werden, man spricht von Enzymhemmung und bezeichnet die Herabsetzung der katalytischen Aktivität eines Enzyms durch einen spezifischen Hemmstoff (Inhibitor, Abb. 42). Es gibt die irreversible Hemmung, bei der ein Hemmstoff eine unter physiologischen Bedingungen nicht umkehrbare Verbindung mit dem Enzym eingeht.

Und es gibt die reversible Hemmung, bei der der gebildete Enzym-Inhibitor-Komplex wieder in seine Bestandteile zerfallen kann. Bei der reversiblen Hemmung unterscheidet man wiederum zwischen

- kompetitiver Hemmung – das Substrat konkurriert mit dem Hemmstoff um die Bindung an das aktive Zentrum des Enzyms. Der Hemmstoff ist aber nicht enzymatisch umsetzbar und stoppt dadurch die Enzymarbeit;

- unkompetitiver Hemmung – der Hemmstoff kann ausschließlich an den Enzym-Substrat-Komplex binden. Er verhindert die katalytische Umsetzung des Substrates zum Produkt – und
- nicht-kompetitiver Hemmung – der Hemmstoff bindet sowohl an das freie Enzym als auch an den Enzym-Substrat-Komplex. Der Enzym-Substrat-Inhibitor-Komplex ist katalytisch inaktiv.

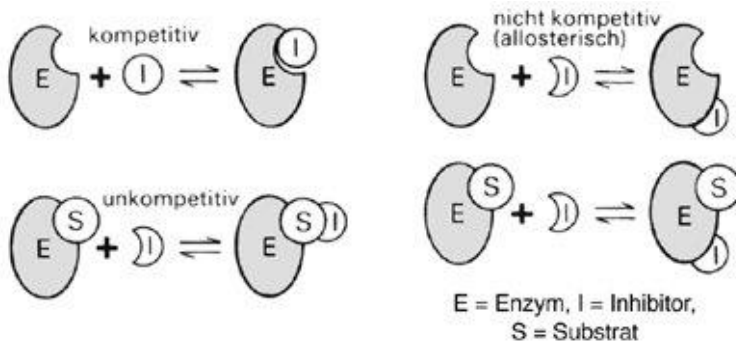


Abbildung 42 Enzymhemmung

ATP

Für die in Zellen ablaufenden Prozesse wird Energie benötigt, da dabei chemische, osmotische oder mechanische Arbeit geleistet wird. Diese Energie muss in irgendeiner Form bereitgestellt werden. Dies geschieht über das Molekül Adenosintriphosphat, kurz ATP. Alle Abbauprozesse des Organismus dienen hauptsächlich der Erzeugung von Energie, die in Form des ATP gespeichert wird. Ein erwachsener braucht täglich etwa 70 kg ATP. Aber die Menge an ATP in unserem Körper beträgt etwa 50g. Das bedeutet letztendlich, dass ein Gramm ATP täglich etwa 1400mal recycelt werden muss. Aber was ist ATP genau?

Das Molekül des Adenosintriphosphats besteht aus einem Adeninrest, dem Zucker Ribose und drei Phosphaten (Abb. 43). Die Bindungen der drei Phosphatreste sind sehr energiereiche chemische Bindungen. Werden diese Bindungen durch Enzyme hydrolytisch gespalten, entsteht das Adenosindiphosphat (ADP) bzw. das Adenosinmonophosphat (AMP). Dabei werden jeweils etwa 32,3 kJ/mol oder 64,6 kJ/mol Energie frei. Diese freiwerdende Energie ermöglicht die Arbeitsleistungen in den Zellen. Aus dem bei der Energieabgabe aus ATP entstandenen AMP bzw. ADP regeneriert die Zelle das ATP. Die Regeneration des ATP geschieht über die Zellatmung.

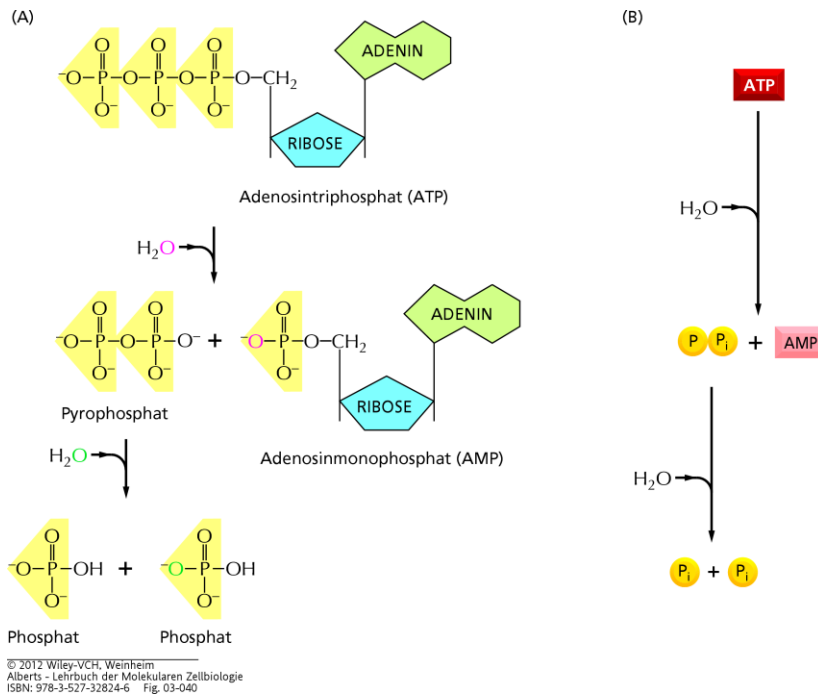


Abbildung 43 ATP

Bei der Zellatmung werden Glukose, Sauerstoff und Wasser in den Mitochondrien zu Wasser und Kohlenstoffdioxid abgebaut. Durch diesen Abbau entsteht Energie in Form von Adenosintriphosphat.

Enzyme spalten dieses ATP in den Mitochondrien in ein Adenosindiphosphat (ADP) und ein freies Phosphat auf. Die dabei freigesetzte Energie wird in Form von Wärme abgegeben, überwiegend jedoch für die Muskelfunktion verwendet.

Anschließend muss das ADP jedoch wieder in ATP umgewandelt werden. So entsteht ein ewiger Kreislauf, der nur bei einer optimalen Versorgung mit Nährstoffen einwandfrei funktioniert.

Wir werden uns im nächsten Kapitel die entscheidenden Schritte der Zellatmung näher anschauen.

Kapitel 5: Biomembranen

Biomembranen bilden eine Barriere und grenzen einerseits die Zelle nach außen ab (=Zellmembran), andererseits umgeben sie die Zellorganellen.

Biomembranen sind eine Doppelschicht aus Phospholipiden, mit deren Aufbau wir uns schon im Beitrag über die Biomoleküle befasst haben (Abb. 36).

Phospholipide haben einen hydrophilen (=wasserliebenden) Kopf und einen hydrophoben (=wassermeidenden) Schwanz. Lagern sich Phospholipide zusammen bilden sie eine Doppelschicht. Die hydrophoben Anteile zeigen nach innen, die hydrophilen Köpfe zeigen nach außen und stehen so in Kontakt mit einem wässrigen Medium, z. B. Wasser oder dem Zellplasma. Diesem Grundaufbau folgen alle Biomembranen. In diese Doppelschicht sind Membranproteine und Kohlenhydrate ein- oder aufgelagert (Abb. 44).

Zu den wichtigsten Aufgaben der Biomembranen gehört einerseits der Stofftransport. Andererseits sind sie in der Lage geschlossene Räume zu bilden, in denen Stoffe gespeichert werden (z.B. in der Vakuole) oder Reaktionen ablaufen können (z.B. in den Mitochondrien).

Flüssig-Mosaik-Modell

Das Flüssig Mosaik Modell erklärt den molekularen Aufbau biologischer Membranen. Es wurde 1972 von Seymour Jonathan Singer und Garth Nicolson entwickelt.

Biomembranen sind laut dem Flüssig-Mosaik-Modell aus einer zweidimensionalen flüssigen Doppelschicht aus Phospholipiden aufgebaut. Die einzelnen Phospholipidmoleküle als auch die darin eingelagerten Membranproteine können sich aber innerhalb dieser Doppelschicht bewegen. Diese Bewegung wird auch als laterale Diffusion bezeichnet. In den meisten Fällen sind die Membranproteine nicht fest mit den Lipidmolekülen der Biomembran verbunden und können so frei innerhalb dieser bewegen. Damit ist die Membranstruktur nicht statisch, sondern dynamisch. Nicht nur die Proteine selbst, sondern auch die Phospholipide können sich seitlich bewegen.

Wie stark diese Bewegung innerhalb der Membran ist, hängt von der Membranfluidität ab. Diese ist wiederum abhängig von der Zusammensetzung der Fettsäuren der Biomembran als auch von der Temperatur. Hier gilt: Je größer die Fluidität, desto dünnflüssiger ist die Membran bzw. je geringer die Fluidität, desto dickflüssiger ist die Membran.

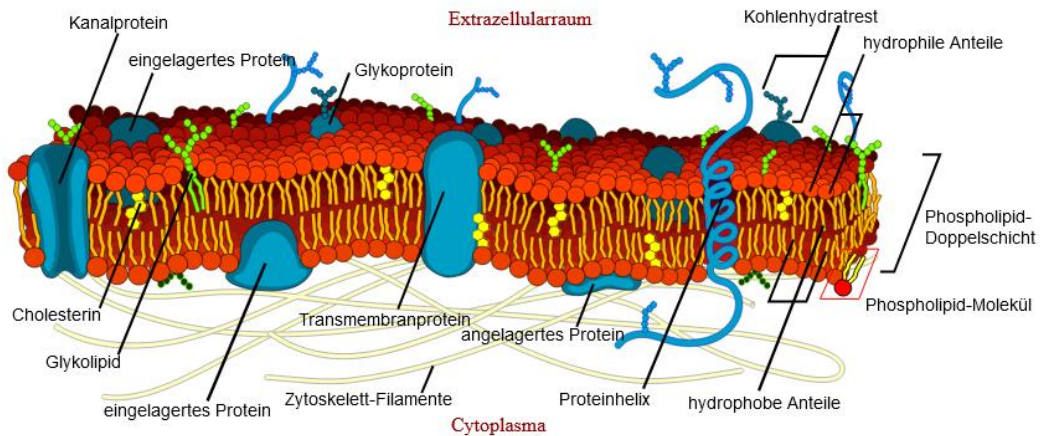


Abbildung 44 Biomembran mit eingelagerten Membranproteinen

Semipermeable Membran und Diffusion

Biomembranen gelten gemein hin als semipermeabel. Das heißt nichts anderes, als dass sie für bestimmte Stoffe durchlässig sind, für andere nicht. Die Durchlässigkeit für die jeweiligen Substanzen hängt vor allem von deren Molekülgröße und Polarität bzw. Ladung ab (Abb. 45).

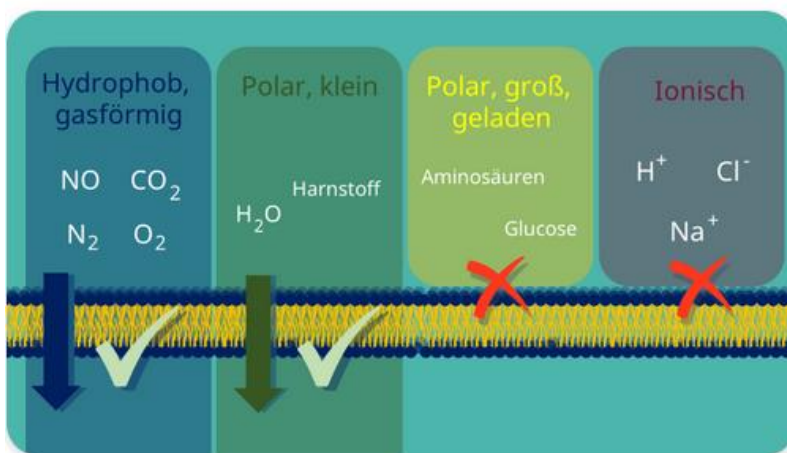


Abbildung 45: Semipermeabilität der Biomembranen.

Grundsätzlich kann man zwischen zwei verschiedenen Prozessen, wie Substanzen semipermeable Membranen passieren können, unterscheiden: dem passiven und dem aktiven Transport.

Kleine hydrophobe und gasförmige Moleküle wie CO_2 und O_2 können ungehindert durch die Biomembran durchdringen. Dasselbe gilt auch für kleine polare Moleküle wie z. B. Wasser. Große Moleküle, wie Aminosäuren und Glucose, sowie Ionen wie H^+ und Salze können die Biomembran nicht durchdringen.

Für den passiven Transport wird keine Energie benötigt, die Moleküle können einfach so durch die Membran durchdringen. Das würde also für CO_2 , O_2 und Wasser gelten.

Diese Bewegung wird als Diffusion bezeichnet und findet immer entlang eines Konzentrationsgefälles statt. Also von einer Region mit einer höheren Konzentration einer Substanz hin zu einer Region mit einer niedrigen Konzentration derselben Substanz. Dies läuft so lange ab, bis sich ein Gleichgewicht eingestellt hat.

Diffusion bewirkt, dass sich gelöste Teilchen in einem Lösungsmittel gleichmäßig und spontan verteilen, bis überall die gleiche Anzahl an Teilchen vorliegt.

Bei der Diffusion handelt sich um einen natürlich ablaufenden, passiven Prozess. Er findet solange statt, bis überall die gleichen Konzentrationen an Teilchen vorliegen. Diffusion als Transportprozess ist aber nicht auf Membranen beschränkt.

Osmose stellt einen Spezialfall der Diffusion dar. Hier diffundiert Wasser in Form von Lösungsmittel entlang seines Konzentrationsgefälles durch eine semipermeable Membran.

Wasser diffundiert aus dem Kompartiment, an dem es höher konzentriert vorliegt (geringere Teilchenkonzentration), in das Kompartiment, in dem seine Konzentration geringer ist (höhere Teilchenkonzentration). Hier bewegen sich also im Gegensatz zur Diffusion nicht die Teilchen, sondern das Lösungsmittel, bis die Stoffkonzentrationen auf beiden Seiten ausgeglichen sind. Es herrscht ein Gleichgewichtszustand vor (Abb. 46).

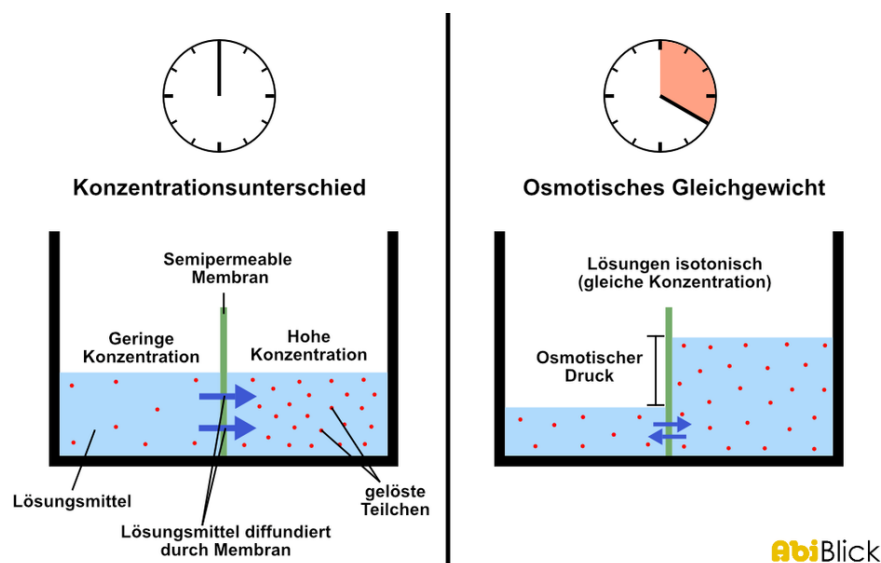


Abbildung 46 Osmose

Der Transport von Wasser zwischen Zellen oder Zellorganellen erfolgt immer vom Ort der höheren zum Ort der niedrigeren Konzentration an gelösten Teilchen. Besonders bei Pflanzenzellen ist der Wassertransport von großer Bedeutung. Aber auch viele Tierzellen sind in einer wässrigen Umgebung von Blut und Lymphe umgeben. Osmose stellt also einen essentiellen Stofftransport im menschlichen Körper dar.

Je größer der Unterschied an gelösten Teilchen auf den jeweiligen Seiten ist, desto größer ist das Bestreben der Seite mit mehr gelösten Teilchen Wasser aufzunehmen, man spricht von osmotischem Druck. Er stellt die Triebkraft dar, damit das Lösungsmittel die Membran passieren kann.

Osmose findet als Transportprozess von Wasser zwischen Zellen statt. Generell gibt es drei Begriffe, um die Konzentrationen von Stoffen, die durch Membranen getrennt sind, zu beschreiben: hypertonisch, isotonisch und hypotonisch.

Eine hypertonische (hyper = über, tonus = Spannung) Lösung besitzt eine höhere Konzentration an gelösten Stoffen als die Vergleichslösung. Bei isotonischen (iso = gleich) Verhältnissen sind auf beiden Seiten gleiche Konzentrationen gelöster Stoffe zu finden. Eine hypotone (hypo = unter) Lösung besitzt eine niedrigere Konzentration an gelösten Stoffen als die Vergleichslösung.

Man kann sich folgenden Zusammenhang merken: Wasser bewegt sich immer von einer hypotonischen Lösung zu einer hypertonischen Lösung.

Am Beispiel tierischer Blutzellen und der Vakuolen der Pflanzen soll dies veranschaulicht werden (Abb. 47).

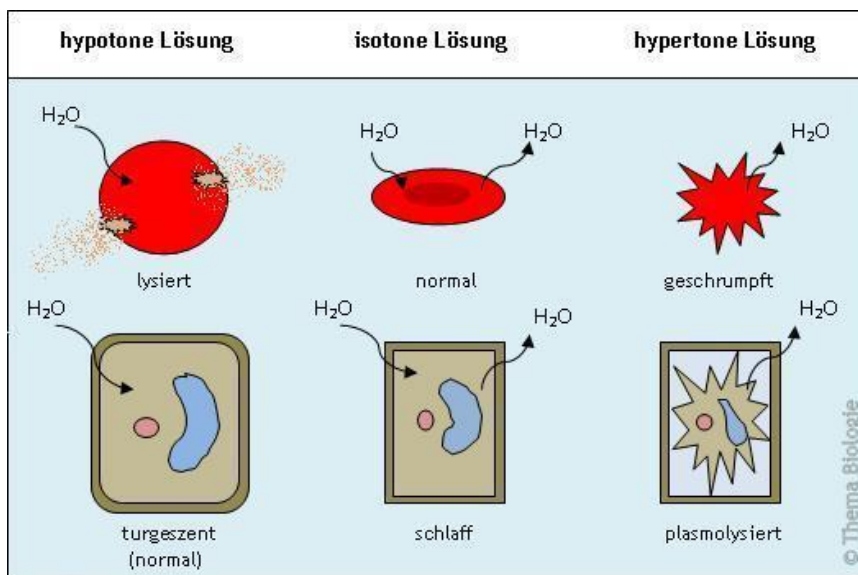


Abbildung 47 Osmose bei Blutzellen (oben) und Vakuolen (unten)

Unser Blut besteht aus Blutplasma und den Blutzellen. Die roten Blutkörperchen sind dabei die häufigsten Blutzellen. Das Blutplasma ist das wässrige Medium des Blutes, indem auch die Salze und Nährstoffe gelöst sind.

Unter normalen (isotonischen) Bedingungen sind rote Blutzellen scheibchenförmig (bikonkav). Verdünnt man jetzt die Blutropfen mit reinem Wasser, schwellen die Zellen an und platzen. Das reine Wasser ist also hypoton im Vergleich zum hypertonen Medium im Inneren der roten Blutzellen. Die Plasmamembran kann dem Druck des Wassereinstroms nicht standhalten und das führt zum Platzen der Zellen. Wird im umgekehrten Fall eine relativ konzentrierte Salzlösung dem Blut zugefügt, ist nun das Zellplasma der Blutzellen hypoton im Vergleich zum Außenmedium. Das hat zur Folge, dass Wasser per Osmose austritt und die Blutzellen eine faltige Form (Stechapfelform) besitzen. Diese sind nun nicht mehr in der Lage, Sauerstoff zu transportieren, was gravierende Folgen hat.

Bei Pflanzen sind die Verhältnisse ähnlich, aber dennoch gibt es einen Unterschied: Pflanzenzellen verfügen über eine feste Zellwand, wodurch sie bei der Aufnahme von zu viel Wasser nicht gleich platzen, denn die Zellwand stabilisiert die Pflanzenzelle. Wasser wird in die Vakuolen aufgenommen, wodurch diese in einem hypotonen Medium (z. B. destilliertes Wasser) anwachsen. Die Vakuole übt dadurch einen starken Druck auf die Zellwand aus. Die Zellwand baut beim Wassereinstrom einen Druck auf – den Turgor. Eine solche prall gefüllte Pflanzenzelle wird auch als turgeszent bezeichnet. Dieser Vorgang wird auch als Deplasmolyse genannt.

Legen wir Pflanzenzellen in eine hypertone Lösung, z. B. in eine Salzlösung, so kommt es zu einem Wasserausstrom aus der Vakuole und der Turgordruck nimmt ab. Der Zellkörper löst sich von der Zellwand. Diesen Vorgang wird auch als Plasmolyse bezeichnet. Für Pflanzen ist der Wassertransport von der Wurzel enorm wichtig. Hier wird Wasser passiv über die Wurzel per Osmose aufgenommen. In biologischen Systemen spielt oft der Begriff Wasserpotential eine Rolle, wenn es um die Beschreibung des Wasserhaushaltes geht. Es beschreibt die Verfügbarkeit von Wasser in einem System wie der Luft oder dem Boden und setzt sich aus dem osmotischen Druck – also der Anzahl der gelösten Teilchen – und dem Turgordruck zusammen. Wichtig: Das Wasser fließt immer vom Ort des höheren Wasserpotentials zum Ort des niedrigeren Wasserpotentials.

Erleichterte Diffusion und Aktiver Stofftransport

Wir wissen nun, dass einige Stoffe die Biomembranen ungehindert diffundieren können. Der passive Transport erfordert keine Energie. Doch was passiert mit den großen Molekülen?

Bei der erleichterten Diffusion können auch polare Stoffe wie Aminosäuren oder Zucker sowie geladene Ionen die semipermeable Membran ohne Verbrauch von Stoffwechselenergie überwinden. Dies ermöglichen sogenannte Membranproteine. Hier gibt es Kanalproteine und Carrier-Proteine (Transportproteine, Abb. 48).

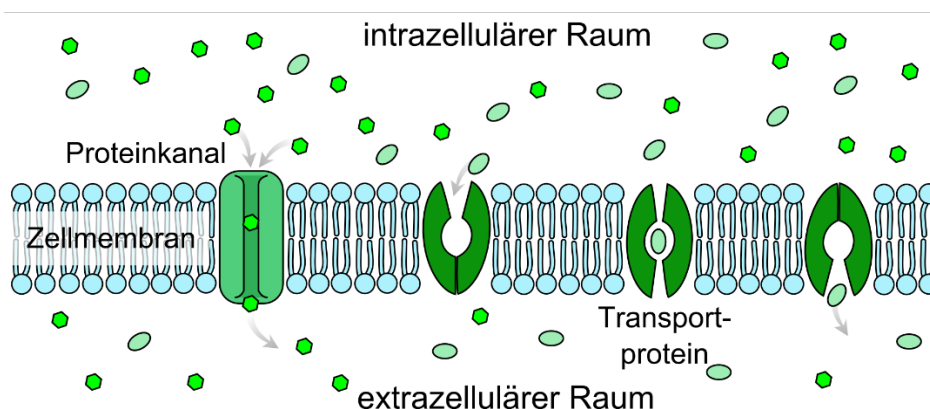


Abbildung 48 Membranproteine

Kanalproteine werden auch als integrale Membranproteine bezeichnet, da sie sich in die hydrophobe Lipidschicht einlagern und so in die Membran integriert sind. Sie bilden

in ihrem Inneren hydrophile Ionenkanäle, die den Durchtritt von Ionen oder anderen größeren Molekülen gewähren. Das Öffnen und Schließen dieser Kanäle wird über Signalproteine oder eine elektrische Spannung reguliert.

Transportproteine hingegen binden an die zu transportierenden Substrate und schleusen sie durch die Doppelschicht der semipermeablen Membran.

Für den Wassertransport besitzt die Zellmembran außerdem die sogenannten Aquaporine. Obwohl Wasser ein sehr kleines Molekül ist und die Membran meist ohne Probleme passieren kann, ist der Wassertransport sehr wichtig. Über die Aquaporine kann die Zelle ihre Aufnahme oder Abgabe von Wasser gezielter regulieren.

Neben dem passiven Stofftransport ist auch ein aktiver Stofftransport möglich. Dieser benötigt, wie der Name bereits vermuten lässt, von außen zugeführte Energie. Mithilfe von integralen Membranproteinen werden Substanzen entgegen ihres Konzentrationsgradienten durch eine Membran transportiert. Es gibt den primär aktiven Transport und den sekundär aktiven Transport.

Beim primär aktiven Transport wird Energie gewonnen, indem ATP mithilfe von Enzymen hydrolysiert wird. Diese Energie treibt den Transport an. Das spielt vor allem in der Neurobiologie bei den Na^+ / K^+ -Ionenpumpen eine große Rolle.

Der sekundär aktive Transport nutzt den Konzentrationsgradienten einer Substanz für den Stofftransport einer anderen Substanz durch eine Membran aus. Dieser Gradient wurde zuvor allerdings über den primären Transport, also die Hydrolyse von ATP-Molekülen, aufgebaut.

Kapitel 6: Zellatmung

In einem der vorherigen Beiträge befassten wir uns mit den Grundlagen der chemischen Reaktion und den Enzymen, sowie dem Adenosintriphosphat. An dieser Stelle wird es Zeit sich mit einigen konkreten chemischen Reaktionen in unserer Zelle zu befassen.

Wir haben gelernt, dass Kohlenhydrate, zu denen u. a. Zucker und Stärke gehören, unsere zentrale Energiequelle sind, auf der fast alle Stoffwechselprodukte beruhen. Doch wie schaffen es unsere Zellen aus diesen Kohlehydraten ihre Energie zu gewinnen?

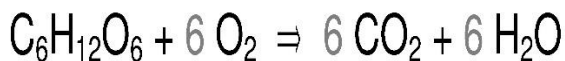
Dies geschieht mittels der Zellatmung.

Wir haben gelernt, chemische Reaktionen durch eine Reaktionsgleichung darzustellen. Wie würde diese Reaktionsgleichung bei der Veratmung von Glucose ausschauen?

Die entsprechende Summenformel der Zellatmung lautet: $C_6H_{12}O_6 + 6 O_2 + 6 H_2O \rightarrow 12 H_2O + 6 CO_2$.

Eine vereinfachte, verkürzte Formel reduziert die Anzahl der Wassermoleküle, da sie sich mathematisch rauskürzen lassen: $C_6H_{12}O_6 + 6 O_2 \rightarrow 6 H_2O + 6 CO_2$.

aerobe / innere Zellatmung



Glukose (Zucker) + Sauerstoff $\Rightarrow$ Kohlendioxid + Wasser

Die **exotherme** Reaktion erzeugt **Energie** als "Zellbrennstoff" **ATP** (Adenosintriphosphat).

Aus einem Molekül Glucose und 6 Molekülen Sauerstoff entstehen 6 Moleküle CO_2 und 6 Moleküle Wasser. Eine chemische Reaktion muss im Gleichgewicht sein, das heißt auf Seiten der Ausgangsstoffe (Edukte) müssen genauso viele Atome zu finden sein, wie auf Seiten der Reaktionsprodukte. Bei der Zellatmung handelt sich um eine exotherme Reaktion, bei der Energie freigesetzt wird, die Reaktionsprodukte (CO_2 und Wasser) haben also ein niedrigeres Energieniveau, als die Edukte (Glukose und Sauerstoff). Die frei gewordene Energie wird in Form von ATP gespeichert.

Bei der Zellatmung werden die organischen Stoffe – hier also die Glukose – oxidiert. Glucose gibt bei dieser Reaktion also Elektronen ab. Der Sauerstoff dient dabei als Oxidationsmittel, der die Elektronen aufnimmt (also einem Elektronenakzeptor entspricht) und zu Wasser reduziert wird. Bei einigen Bakterien kommt außerdem die anaerobe Atmung vor, bei der anstelle von Sauerstoff andere Stoffe, z. B. Nitrat (NO_3^-)

) bei der Nitratatmung oder Sulfat SO_4^{2-} bei der Sulfatatmung, als Oxidationsmittel dienen.

Von der Atmung ist die Gärung zu unterscheiden. Hierbei werden Moleküle zur Energiegewinnung abgebaut ohne Einbeziehung externer Elektronenakzeptoren wie Sauerstoff (O_2) oder Nitrat (NO_3^-)

Werden Moleküle abgebaut, um Energie zu erzeugen, spricht man von Katabolismus, während beim Anabolismus Moleküle unter Energieverbrauch synthetisiert werden.

Nun sollte man sich aber nicht vorstellen, dass man einfach nur ein Stück Zucker in die Zelle einzufügen braucht und wie durch ein Wunder entsteht dabei die Energie, die wir brauchen. Damit aus Glucose Energie gewonnen werden kann, braucht es mehrere Zwischenschritte, die durch Enzyme katalysiert werden.

Wir werden uns daher vornehmlich mit dem beschäftigen, was auf dem Reaktionspfeil geschieht.

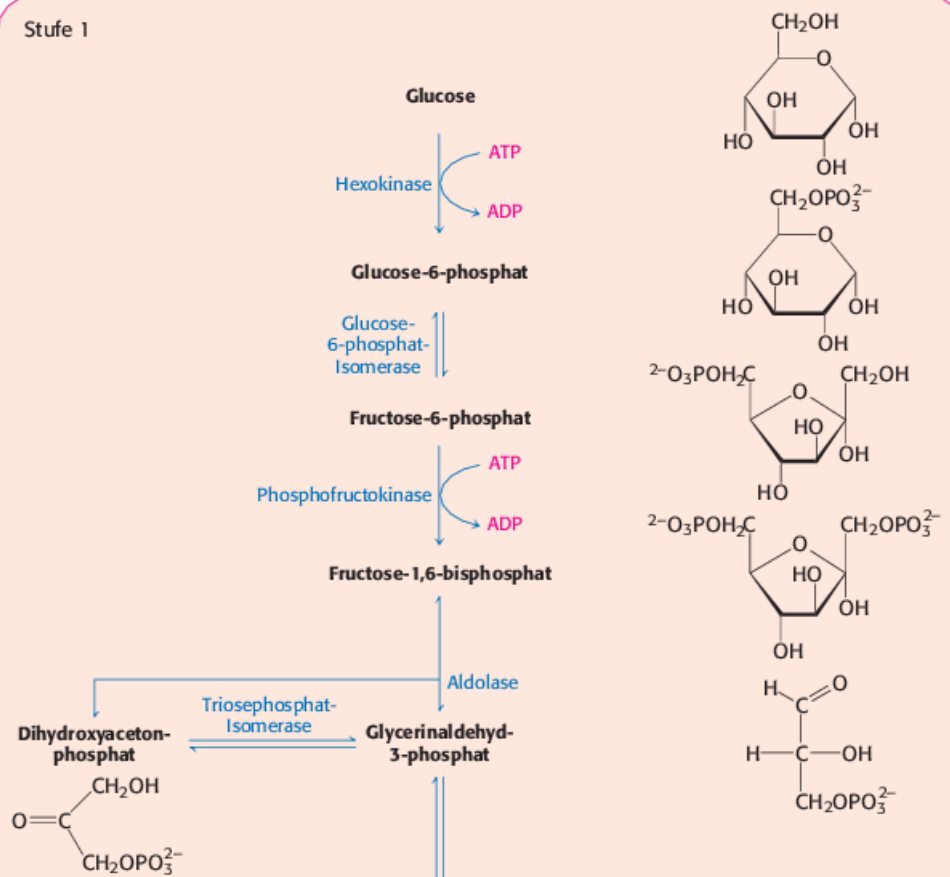
Die Zellatmung umfasst folgende Teilprozesse:

1. die Glykolyse,
2. den Citratzyklus und
3. die Endoxidation in der Atmungskette.

Glykolyse

Die Glykolyse ist der erste Teilschritt der Zellatmung und findet im Zellplasma statt. Dieser Prozess kann ohne Sauerstoff ablaufen, ist also anaerob. Hierbei wird aus einem Molekül Glucose zwei Moleküle Pyruvat hergestellt. Pyruvat ist das Anion (negativ geladenes Atom) der Brenztraubensäure und ist die Ausgangssubstanz für viele weitere Stoffwechselwege. Die Bildung von Pyruvat aus Glucose erfolgt in 10 Teilschritten (Abb. 49), die in zwei Phase zusammengefasst werden. Bei Phase 1 wird Energie in Form von ATP verbraucht, bei Phase 2 ATP gebildet. Wir schauen uns diese Prozesse näher an.

Stufe 1



Stufe 2

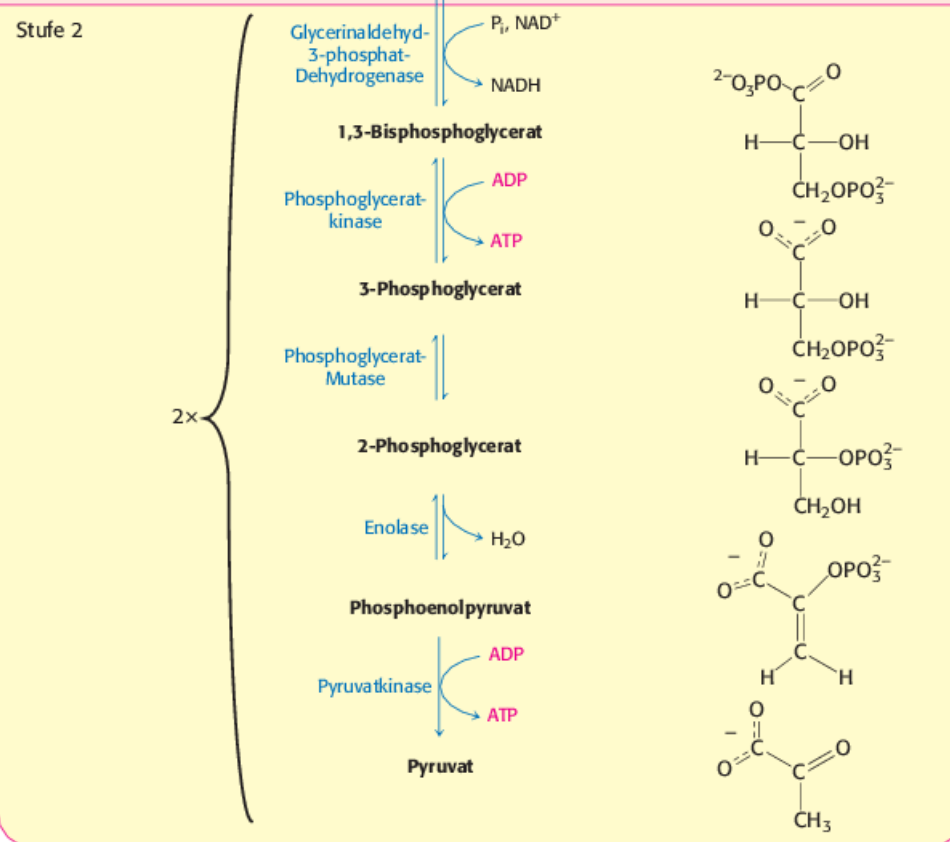


Abbildung 49 Glykolyse

Phase 1.1.: Gelangt Zucker ins Zellplasma wird dieses mittels des Enzyms Hexokinase phosphoryliert und es entsteht Glucose-6-Phosphat. Phosphorylierung bedeutet, dass einem Molekül eine Phosphatgruppe angeheftet wird. Hierbei wird ein Molekül ATP verbraucht, welches dann zum ADP (Adenosindiphosphat) wird. Dieser Erste Schritt kostet also Energie, bezahlt wird mit einem ATP.

Phase 1.2.: Glucose-6-Phosphat wird dann in Fructose 6-Phosphat umgewandelt, dies wird durch das Enzym Glucose-6-Phosphat-Isomerase katalysiert.

Phase 1.3.: Der dritte Schritt kostet wieder Energie, es wird also mit einem Molekül ATP bezahlt: Aus Fructose-6-Phosphat wird Fructose-1,6-bisphosphat. Das Enzym Phosphofructokinase katalysiert diesen Schritt. Wie der Name Fructose-1,6-bisphosphat andeutet, hat dieses Molekül zwei Phosphatgruppen angeheftet, eines am 1. Kohlenstoffatom und die zweite am 6. Kohlenstoffatom.

Phase 1.4.: Mittels des Enzyms Aldolase wird die Fructose-1,6-bisphosphat gespalten in zwei Molekülen mit je 3 Kohlenstoffatomen, sogenannte Triosephosphate. Einmal entsteht Glycerinaldehyd-3-phosphat und Dihydroxyacetonphosphat.

Phase 1.5.: Mittels des Enzyms Triosephosphat-Isomerase kann aus dem Dihydroxyacetonphosphat das Glycerinaldehyd-3-Phosphat gebildet werden. dieses gilt dann als Ausgangsstoff für die Phase 2 der Glykolyse.

Fassen wir eine Zwischenbilanz: in Phase 1 der Glykolyse wird aus einem Molekül Glukose zwei Moleküle Glycerinaldehyd-3-Phosphat gebildet. Diese Umwandlung kostet aber Energie in Form von zwei ATP. D. h. die Energiebilanz ist erstmal negativ. Die Energiegewinnung erfolgt über die Phase zwei. Hier sei erwähnt, dass Phase zwei doppelt abläuft. Wir haben ja nun zwei Moleküle Glycerinaldehyd-3-Phosphat

Phase 2.1.: Glycerinaldehyd-3-Phosphat wird über das Enzym Glycerinaldehyd-3-Phosphat-Dehydrogenase zu 1,3-Bisphosphoglycerat gebildet. Es kommt also eine zweite Phosphatgruppe hinzu. Diese Reaktion ist energetisch sehr günstig und hierfür wird anorganisches Phosphat genutzt (hier als P_i dargestellt), also welches, dass nicht an ATP gebunden ist. Bei dieser Reaktion tritt ein weiteres Molekül auf, dass bei chemischen Reaktionen in der Zelle äußerst wichtig sein kann: das NAD^+ . NAD steht für Nicotinamid-Adenin-Dinukleotid und ist ein Coenzym und Oxidationsmittel und somit an zahlreichen Redoxreaktionen des Stoffwechsels der Zelle beteiligt. Es nimmt also Elektronen und wird zu $NADH+H^+$ reduziert.

Phase 2.2.: 1,3-Bisphosphoglycerat hat nun zwei Phosphatgruppen. Mittels des Enzyms Phosphatglyceratkinase wird eine dieser Phosphatgruppen auf ein Molekül ADP übertragen und es entsteht ATP. Als Produkt entsteht 3-Phosphoglycerat. Da dieser Prozess ja doppelt abläuft entstehen zwei Moleküle ATP, womit die verbrauchten zwei ATP-Moleküle aus Phase 1 wieder „abbezahlt“ wurden. Die Energiebilanz liegt nun bei null.

Phase 2.3.: Aus 3-Phosphoglycerat wird über das Enzym Phosphoglycerat-Mutase 2-Phosphoglycerat. Bei diesem Schritt handelt es sich um eine einfache Umlagerung der Phosphatgruppe vom C3 zum C2-Atom.

Phase 2.4.: Aus 2-Phosphoglycerat wird über das Enzym Enolase Phosphoenolpyruvat. Dabei wird Wasser abgespalten.

Phase 2.5.: Im letzten Schritt der Glykolyse wird aus dem Phosphoenolpyruvat über das Enzym Pyruvatkinase die Phosphatgruppe auf ein Molekül ADP übertragen und es entsteht ATP. Als Produkt haben wir nun endlich das Pyruvat. Da auch dieser Prozess zwei Mal stattfindet entstehen zwei Moleküle ATP und damit haben wir eine positive Energiebilanz. Außerdem sind während der Phase 2 noch zwei Moleküle $\text{NADH} + \text{H}^+$ entstanden, die später wichtig sein werden. Außerdem wurde pro Molekül Glucose 2 Moleküle Wasser abgespalten.

Fassen wir die Bilanz der Glykolyse zusammen:



Gärung und Gluconeogenese

Das Endprodukt der Glykolyse ist Pyruvat, das aufgrund seiner vielseitigen Verwendbarkeit im Stoffwechsel eine zentrale Position einnimmt. Steht kein Sauerstoff zur Verfügung, kann Pyruvat in verschiedenen Gärungen verwendet werden. Am bekanntesten ist die alkoholische Gärung, wozu manche Mikroorganismen, z. B. Hefe, in der Lage sind bei dieser Reaktion wird keine Energie gewonnen, aber das Koenzym NAD^+ , der bei der Glykolyse verbraucht wurde, regeneriert (Abb. 50).

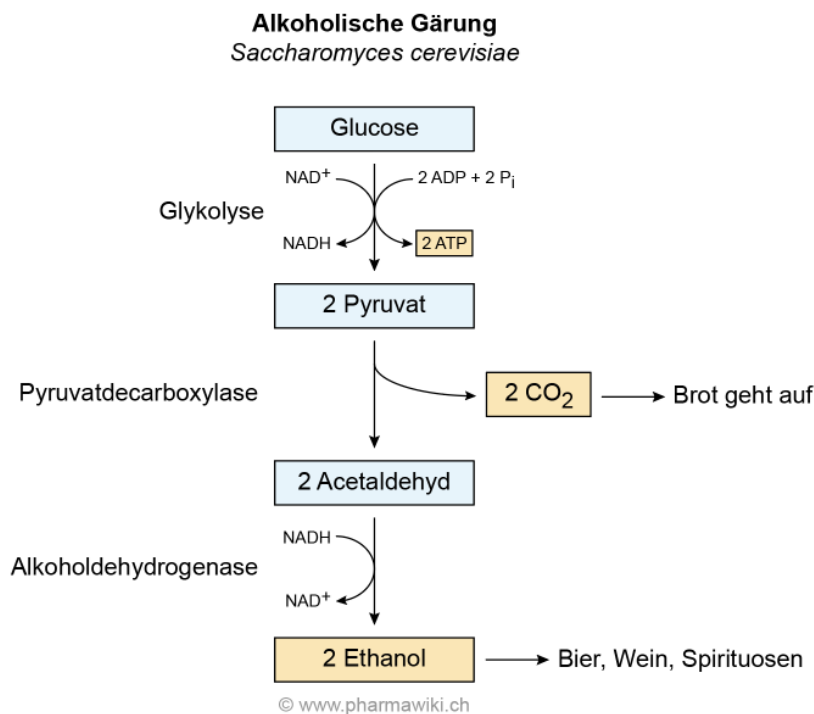


Abbildung 50 Alkoholische Gärung

Die alkoholische Gärung verläuft in zwei Schritten. Im ersten Schritt wandelt die Pyruvatdecarboxylase das Pyruvat in Acetaldehyd um. Dabei wird CO_2 abgespalten. Jeder der einen Hefeteig benutzt hat weiß das, denn durch diese Reaktion geht der Teig auf.

In Schritt zwei wird das Acetaldehyd über die Alkoholdehydrogenase in Ethanol umgewandelt und das ist der Alkohol der in unserem Bier, Wein oder Wodka drin ist. Dabei wird das $\text{NADH} + \text{H}^+$ zu NAD^+ regeneriert.

Andere Mikroorganismen nutzen andere Formen der Gärung. So können einige Bakterien aus Pyruvat Milchsäure herstellen. Dieser als Lactatgärung bezeichnete Prozess findet auch in unseren Muskeln statt, wenn beim Ausdauersport unsere Muskeln nicht mit genügend Sauerstoff versorgt werden.

Übrigens kann auch über die Gluconeogenese aus Pyruvat, aber auch aus anderen Molekülen wie Milchsäure, wieder Glucose hergestellt werden. Dieser Reaktionsweg wird immer dann eingeschlagen, wenn der Nachschub an Kohlehydraten gering ist (z. B. wenn man hungert). In manchen Fällen ist die Gluconeogenese eine Umkehrung der Glykolyse, weil z. B. dieselben Zwischenprodukte vorkommen und auch einige Enzyme in der Lage sind die Reaktion in beide Richtungen ablaufen zu lassen. Aber Glykolyse und Gluconeogenese unterscheiden sich durch vier Enzyme, sodass jeder Stoffwechselweg für sich reguliert wird und sich die beiden Reaktionswege nicht gegenseitig behindern. Des Weiteren findet die Gluconeogenese nicht im Zellplasma sondern in den Mitochondrien statt.

Mittels der Glykolyse kann Energie in Form von ATP gewonnen werden. Aber die Energiebilanz ist eher dürftig, denn unterm Strich bleiben zwei Moleküle ATP übrig. Beim Vorhandensein von Sauerstoff kann aber das Pyruvat weiter oxidiert werden und mehr ATP-Moleküle gewonnen werden. Der nächste Schritt ist nämlich der Citratzyklus.

Bei Citratzyklus, auch Krebszyklus oder Tricarbonsäurezyklus genannt, wird das Pyruvat weiter oxidiert. Der Citratzyklus findet nicht im Zellplasma, sondern in den Mitochondrien statt. Das Ausgangsmolekül dieses Reaktionsweges ist das Acetyl-Coenzym A, kurz AcetylCoA genannt. Dessen Bildung schauen wir uns genauer an

Oxidative Decarboxylierung von Pyruvat

Die Bildung von Acetyl-CoA ist ein mehrstufiger Prozess, welches in den Mitochondrien stattfindet und durch ein Multienzymkomplex katalysiert. Der Name dieses Multienzymkomplexes ist Pyruvatdehydrogenasekomplex (PDH-Komplex) und besteht aus drei Enzymen, die die wichtigen Teilschritte der Acetyl-CoA-Bildung katalysieren. Außerdem sind 5 Koenzyme beteiligt: Thiaminpyrophosphat (TPP), Liponsäure, FAD, Coenzym A und NAD^+ .

Es ist also eine recht komplexe Reaktion mit mehreren Zwischenschritten. Für uns wichtig ist aber folgendes (Abb.51):

Vom Pyruvat wird CO_2 abgespalten und ein Acetyl-Rest entsteht. Der C3-Körper des Pyruvats wird also zum C2-Körper, weil nun ein CO_2 abgespalten wurde. Dieser Acetyl-Rest wird oxidiert und es NAD^+ wird zu NADH reduziert. Die Acetylgruppe wird auf das Coenzym A übertragen, sodass ein Molekül Acetyl-CoA entsteht.

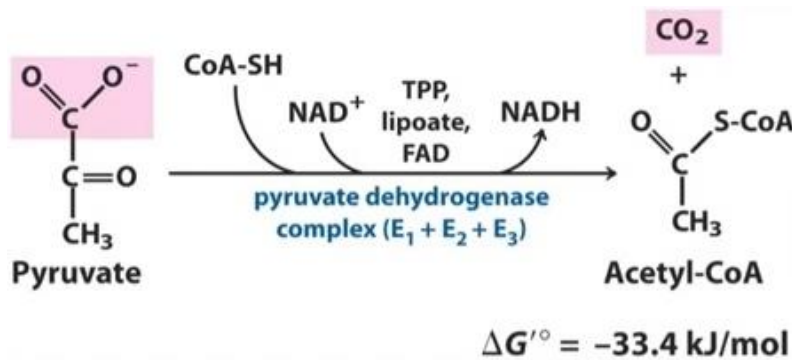


Abbildung 51 Bildung von Acetyl-CoA

Acetyl-CoA dient als Brennstoff für den Citratzyklus in der nächsten Phase der Zellatmung. Die Bindung von CoA hilft dabei, die Acetylgruppe zu aktivieren und bereitet sie auf die notwendigen Reaktionen vor, um in den Citratzyklus einzutreten.

Citratzyklus

Der Citratzyklus ist ein zyklischer Stoffwechselprozess, der in den Mitochondrien stattfindet. Er wird gerne als die Drehscheibe des Stoffwechsels bezeichnet, da er eine zentrale Rolle für viele Stoffwechselwege einnimmt. So sind einige seiner Zwischenprodukte Ausgangsstoffe für die Synthese von Aminosäuren oder Fettsäuren. Sie werden dafür vom Citratzyklus für die entsprechenden Stoffwechselwege abgezweigt. Die wichtigste Funktion des Citratzyklus ist jedoch die Gewinnung von Elektronen (vor allem NADH) für die Atmungskette, also dem dritten Schritt der Zellatmung, durch das Oxidieren von Acetyl-CoA. Die einzelnen Reaktionsschritte des Citratzyklus werden in Abb. 52 dargestellt.

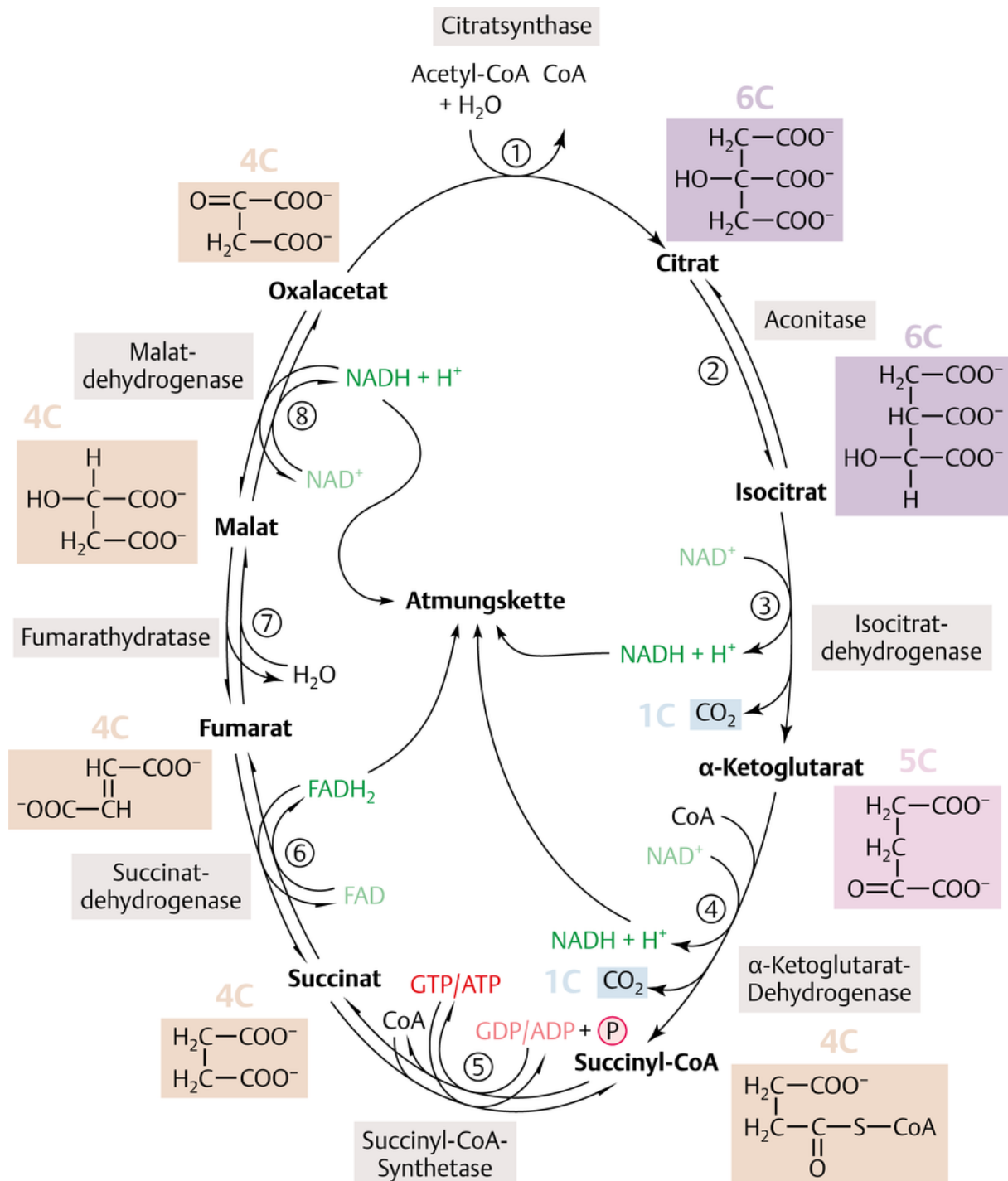


Abbildung 52 Citrat-Zyklus

Hier finden sich 8 Schritte, die wir uns näher anschauen.

Schritt 1: Das Enzym Citratsynthase katalysiert die Übertragung von Acetyl-CoA auf Oxalacetat unter der Bildung von Citrat. Oxalacetat ist ein Molekül mit 4 Kohlenstoffatomen, Acetyl-CoA mit 2. Citrat ist somit ein Molekül mit 6 Kohlenstoffatomen. Bei dieser Reaktion wird H₂O eingeführt und das Coenzym A abgespalten.

Schritt 2: Citrat wird durch das Enzym Aconitase zu Isocitrat umgewandelt. Dabei wird die OH-Gruppe des Citrats verschoben. Diese Reaktion kann problemlos in beide Richtungen verlaufen.

Schritt 3: Die Isocitratdehydrogenase katalysiert das Isocitrat zu α -Ketoglutarat. Dabei wird ein NAD^+ zu $\text{NADH} + \text{H}^+$ reduziert und es spaltet sich ein CO_2 -Molekül ab. In diesem Reaktionsschritt findet also die erste Oxidationsreaktion sowie die erste Decarboxylierung des Citratzyklus statt – mit Bildung von 1 $\text{NADH} + \text{H}^+$ und Freisetzung von CO_2 .

Schritt 4: Die α -Ketoglutarat-Dehydrogenase ist ein großer Enzymkomplex, welcher sehr der Pyruvatdehydrogenase bei der Acetyl-CoA-Synthese ähnelt. Er benötigt für die oxidative Decarboxylierung von α -Ketoglutarat zu Succinyl-CoA folgende Cofaktoren: Thiaminpyrophosphat, Liponamid, Coenzym A, FAD und NAD^+ . Es entsteht erneut CO_2 sowie 1 weiteres $\text{NADH} + \text{H}^+$ für die Atmungskette.

Schritt 5: Aus Succinyl-CoA wird mittels des Enzyms Succinyl-CoA-Synthetase Succinat. Das Coenzym A wird abgespalten. Die dabei freiwerdende Energie wird genutzt um 1 GTP aus GDP zu synthetisieren – auch Substratkettenphosphorylierung genannt. GTP steht für Guanosintriphosphat und ist mit dem ATP verwandt. Überträgt man eine Phosphatgruppe des GTP auf ADP erhält man ATP: $\text{GTP} + \text{ADP} \rightarrow \text{GDP} + \text{ATP}$. Diese Reaktion ist selbst jedoch kein Teil des Citratzyklus.

Schritt 6: Die FAD-abhängige Succinat-Dehydrogenase führt die Oxidation von Succinat zu Fumarat durch. Dies geschieht unter der Ausbildung einer Doppelbindung und Freisetzung von 1 FADH_2 . FAD steht für Flavin-Adenin-di-Nukleotid und ist ein Coenzym vergleichbar mit NAD^+ .

Schritt 7: Die Anlagerung von Wasser an Fumarat katalysiert die Fumarat-Hydratase – auch Fumarase genannt – und führt zur Ausbildung von Malat.

Schritt 8: Die NAD^+ -abhängige Malat-Dehydrogenase oxidiert Malat zu Oxalacetat, welches dann erneut als Substrat für Schritt 1 des Citratzyklus vorliegt. Hierbei entsteht 1 $\text{NADH} + \text{H}^+$ für die Atmungskette.

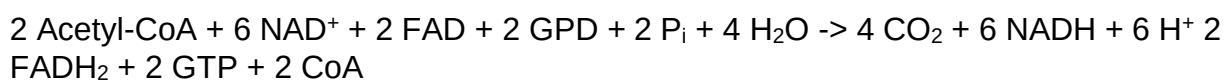
Bilanz:

Bei der oxidativen Decarboxylierung dem dem Citratzyklus können folgende Bilanzen ermittelt werden:

Bilanz der oxidativen Decarboxylierung von Pyruvat:



Bilanz des Citratzyklus:



Das entstandene CO_2 ist unter diesem Gesichtspunkt lediglich ein Abfallprodukt zu verstehen. Für den nächsten Schritt, die Atmungskette, sind die 10 Moleküle $\text{NADH}+\text{H}^+$ und 2 FADH_2 Moleküle entscheidend.

Atmungskette

Im Citratzyklus werden überwiegend $\text{NADH}+\text{H}^+$ und FADH_2 produziert, die den Brennstoff für die Atmungskette und damit für die Synthese von ATP liefern. Man spricht auch von oxydativer Phosphorylierung. Die beiden reduzierten Coenzyme transportieren Elektronen, die bei der Oxidation von Pyruvat angefallen sind. Diese Elektronen werden für die Reduktion von Sauerstoff zu Wasser benötigt.

Hier heißt es durchatmen, denn die Atmungskette ist etwas schwerer nachzuvollziehen als die Glykolyse oder der Citratzyklus.

Die Atmungskette kann als Elektronentransportkette verstanden werden. Elektronentransportketten bestehen aus einer Reihe hintereinander geschalteter Redox-Moleküle, die in der Lage sind, Elektronen aufzunehmen bzw. abzugeben. Über diese Kette werden Elektronen von höheren Energieniveaus auf niedrigere weitergegeben, sie fallen sozusagen in Stufen bergab, wobei die einzelnen Redox-Moleküle ein zunehmend niedriges Energieniveau haben.

Die Enzyme der Atmungskette finden sich bei Eukaryoten (Zellen mit Zellkern) in der inneren Mitochondrienmembran, bei Prokaryoten (Zellen ohne Zellkern, wie Bakterien) in der Plasmamembran. Hier laufen zahlreiche Redoxreaktionen ab, also Reaktionen bei denen die Oxidation (Abgabe von Elektronen) einer Verbindung mit der Reduktion (Aufnahme von Elektronen) einer anderen Verbindung gekoppelt ist. In der Atmungskette werden die Elektronen von $\text{NADH}+\text{H}^+$ und FADH_2 schrittweise auf Sauerstoff übertragen, welcher dann zu Wasser reduziert wird.

Ein Mitochondrium ist ein von einer Doppelmembran umschlossenes Zellorganell, das eukaryotischen Zellen (d.h. Zellen, die einen Zellkern haben) zur Energiegewinnung dient (Abb. 53).

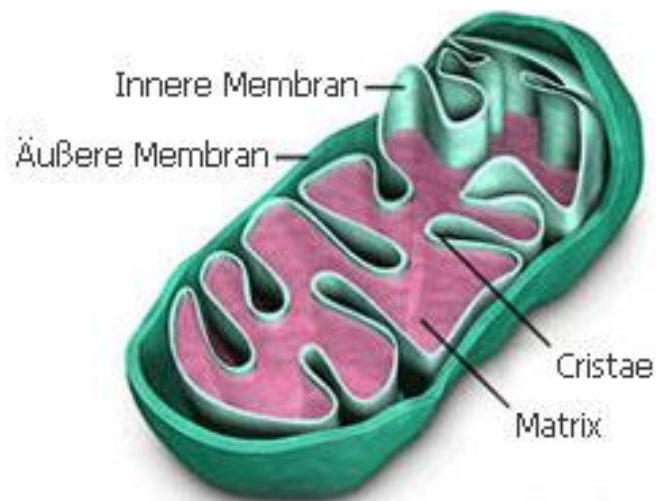


Abbildung 53 Mitochondrium

Die äußere Membran grenzt das Mitochondrium nach außen ab und enthält Kanäle für die Durchlässigkeit von Molekülen. Die innere Membran bildet große Einfaltungen, die zisternenförmig das Innere des Mitochondriums ausfüllen (Cristae). In der Membran der Cristae befinden sich die Enzymkomplexe der Atmungskette und die ATP-Synthase, ein Transmembranprotein, dessen Funktion die Produktion von ATP ist.

Im inneren Membranraum findet sich die mitochondriale Matrix. In der Matrix findet man die ringförmige DNA des Mitochondriums und Ribosomen. Alle wesentlichen Enzyme des Citratzyklus sind im Matrixraum der Mitochondrien und daher in direkter Nachbarschaft der Atmungskette lokalisiert.

Abb. 54 zeigt die innere Mitochondrienmembran und ihre für die Atmungskette entscheidenden Bereiche.

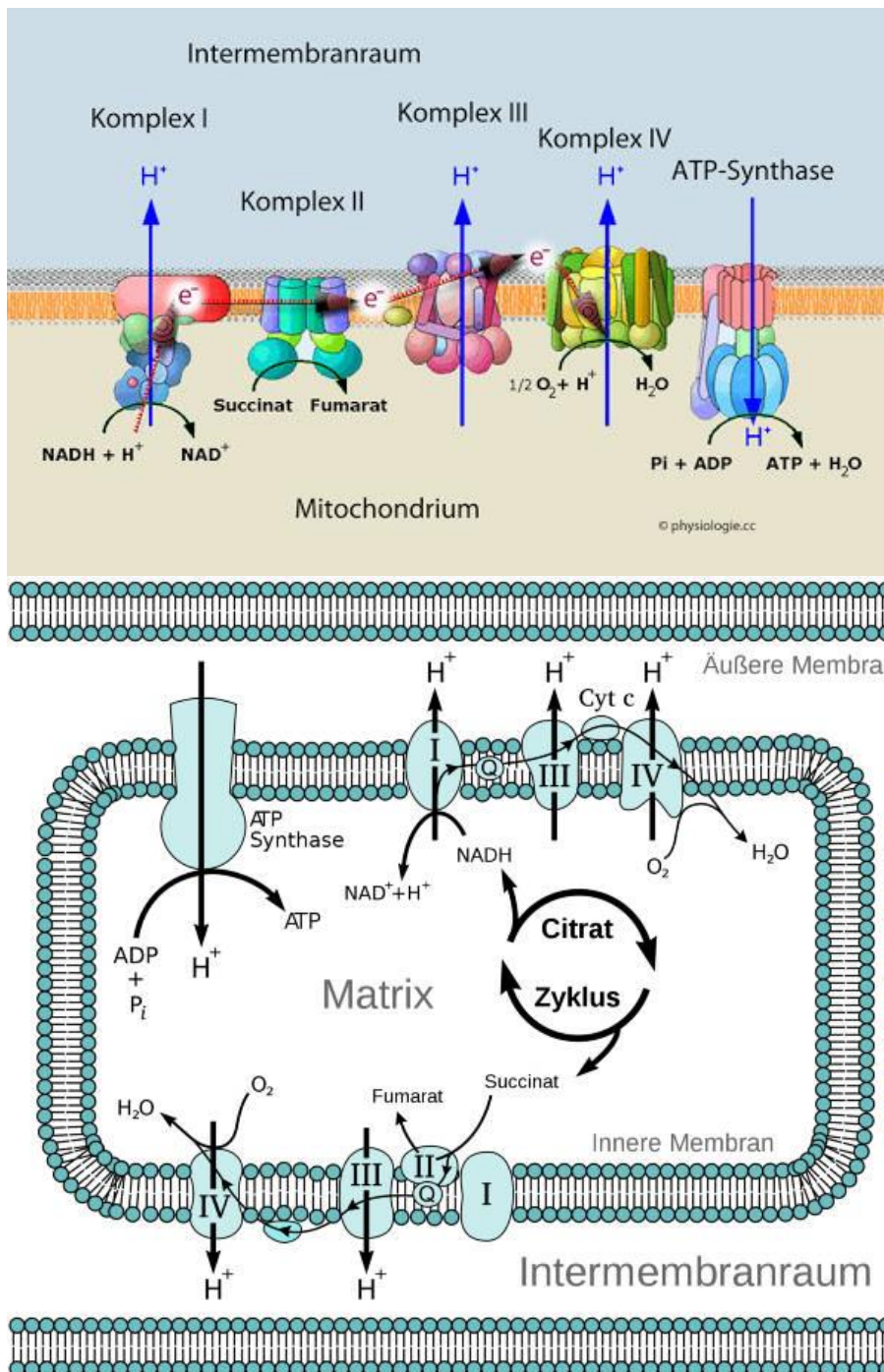


Abbildung 54 Atmungskette

Die mitochondriale Atmungskette besteht aus einer Reihe von Proteinkomplexen. Diese wirken als Oxireduktasen und werden als Komplex I-IV bezeichnet. Zusätzlich kommt die ATP-Synthase und die Wasserstoff- bzw. Elektronenüberträger Ubichinon (Coenzym Q) und Cytochrom c, die in die innere Mitochondrienmembran eingelagert sind, hinzu.

Die vier Komplexe der Atmungskette bestehen aus einer Reihe von Proteinen und enthalten Reaktionszentren mit Flavinen, Eisenschwefelkomplexen und Cytochrome, die die Übertragung der Wasserstoffe und Elektronen ermöglichen. Neben diesen dienen die Redoxhilfssubstrate Cytochrom und Ubichinon als Sammelbecken für

Elektronen bzw. Wasserstoff. Zwischen den Proteinkomplexen finden Elektronenübergänge statt, die dabei freiwerdende Energie wird sofort zur Erzeugung bzw. Aufrechterhaltung des Protonengradienten umgesetzt.

Die aus dem Citratzyklus stammenden Coenzyme $\text{NADH} + \text{H}^+$ und FADH_2 geben Elektronen ab, die Protonen sammeln sich im Intermembranraum.

Die vier Komplexe der Atmungskette sind:

Komplex I: NADH-Dehydrogenase

Komplex II: Succinatdehydrogenase

Komplex III: Cytochrom-c-Reduktase

Komplex IV: Cytochrom-c-Oxidase

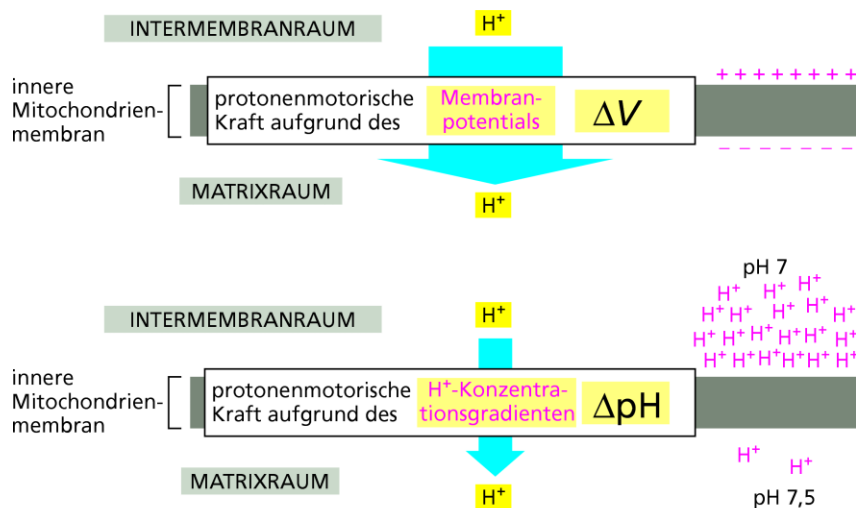
Über die Komplexe I, III und IV werden Protonen über die Membran gepumpt. Die Mehrzahl der Enzyme in diesen Komplexen, die die Elektronen weiter transportieren, verfügen über Eisenionen, die bei Elektronenaufnahme reduziert und bei Elektronenabgabe oxidiert werden. Da die Enzyme eine unterschiedliche Elektronegativität besitzen, werden die Elektronen von Enzym zu Enzym weitergegeben. Das nächstfolgende Enzym hat eine höhere Elektronegativität als das vorhergegangene. Es nimmt deshalb die Elektronen auf. Dadurch kann die Elektronentransportkette stattfinden.

Die Elektronentransportkette beginnt in der Mitochondrienmatrix mit $\text{NADH} + \text{H}^+$, das Elektronen an Komplex I abgibt, also zu NAD^+ oxidiert. Die Elektronen gelangen über Ubichinon zu Komplex III. Ubichinon ist in der Membran beweglich und steht in Kontakt mit Komplex II, der Elektronen von FADH_2 erhält. Komplex III gibt Elektronen an Cytochrom c, dieses bindet an Komplex IV. Komplex IV reduziert letztlich den elementaren Sauerstoff, sodass Wasser entsteht 2 Elektronen werden auf ein halbes O_2 -Molekül übertragen, wobei Wasser (H_2O) entsteht. Für die Arbeit des Komplexes IV ist Sauerstoff als finaler Elektronenakzeptor unersetzlich. Dies ist der Hauptgrund, warum Lebewesen Sauerstoff zum Leben brauchen.

Der Komplex II, die Succinatdehydrogenase, ist direkt an einer Reaktion des Citratzyklus beteiligt und schleust seine Elektronen von dort über enzymgebundenes FAD in die Atmungskette ein, übernimmt in der Atmungskette also keine Protonentransportfunktion.

Die Energie für die ATP-Synthese liegt im von der Elektronentransportkette erzeugten Protonen-Konzentrationsgefälle (sog. Chemiosmose), es entsteht also ein Konzentrationsunterschied (=chemisches Potential), in diesem Fall ein Protonengradient, da wir es mit H^+ mit Protonen zu tun haben.

Die H^+ -Konzentrationen im Intermembranraum (hohe H^+ -Konzentration) und in der Mitochondrienmatrix (niedrige H^+ -Konzentration) unterscheiden sich stark. Damit haben wir auch einen unterschiedlichen pH-Wert. Zusätzlich entsteht ein Spannungsunterschied (elektrisches Potential), da durch die Verteilung der H^+ -Ionen der Intermembranraum positiver ist als die Mitochondrienmatrix (Abb. 55).



© 2012 Wiley-VCH, Weinheim
 Alberts - Lehrbuch der Molekularen Zellbiologie
 ISBN: 978-3-527-32824-6 Fig. 14-0010

Abbildung 55 elektrochemisches Potential

Der Gradient wird durch den Export von H^+ von Komplex I, III und IV aufgebaut.

Biomembranen wie die Mitochondrienmembran bestehen aber aus einer Doppelschicht aus Phospholipiden und bilden dadurch eine Barriere für die geladenen Wasserstoffprotonen. Die H^+ -Ionen sind daher gefangen und können nicht frei durch die Membran diffundieren. Sie können nur durch ein Kanalprotein – die ATP-Synthase – diffundieren, um den Konzentrations- und Ladungsunterschied auszugleichen.

Das Enzym ATP-Synthase nutzt die Energie aus dem Gradienten und katalysiert beim Rücktransport der Protonen die ATP-Synthese. Das Enzym besteht aus einem Protonenkanal und einer zur Matrix gerichteten Kopfstruktur.

Der Protonenrückstrom erzeugt Energie. Das kann man sich wie bei einem Wasserkraftwerk vorstellen, bei dem Wasser hinter einer Staumauer angestaut wird. Dieser Rückstau wird verwendet, um die Turbinen zur Rotation zu bringen, um Strom zu generieren. Diese Turbine ist unserem Fall die ATP-Synthase. Sie koppelt die Diffusion der Protonen, wie ihr Name vermuten lässt, mit der Synthese von ATP aus ADP und einer Phosphatgruppe.

Nach Bindung von H^+ findet eine Konformationsänderung statt, sodass das Proton in die Matrix abgegeben werden kann. Dabei rotiert die Kopfstruktur wie ein Rotationsmotor und erzeugt ATP-Moleküle.

Endbilanz

Glucose wurde somit zu den energiearmen Produkten CO_2 und H_2O abgebaut. Doch wie viel ATP entsteht bei der Atmungskette?

Bei der Oxidation eines $NADH+H^+$ Moleküls aus dem Citratzyklus entstehen 2,5 Moleküle ATP. Die $NADH+H^+$ -Moleküle aus der Glykolyse befinden sich im Cytoplasma und müssen erstmal zu den Mitochondrien gebracht werden. Dies geschieht mittels Membranproteinen, die dieses $NADH+H^+$ in den Intermembranraum der Mitochondrien transportieren, so z. B. das Glycerin-3-phosphat-Shuttle. Dadurch

ist die Ausbeute von ATP in diesem Fall geringer, etwa 1,5 Moleküle ATP pro $\text{NADH} + \text{H}^+$ aus der Glykolyse.

Wird jedoch das $\text{NADH} + \text{H}^+$ der Glykolyse durch den Malat-Aspartat-Shuttle, einem anderen Membranprotein, in die Matrix gebracht, können daraus 2,5 ATP erzeugt werden. Damit können maximal $10 \times 2,5 = 25$ ATP (2 $\text{NADH} + \text{H}^+$ aus Glykolyse + 8 aus dem Citratzyklus) erzeugt werden.

Mit dem FADH_2 verläuft der Vorgang im Prinzip genauso, nur gibt FADH_2 auf einem niedrigeren Energieniveau Elektronen ab, sodass pro FADH_2 maximal 1,5 ATP entstehen können. Da zwei FADH_2 oxidiert werden, entstehen dabei 3 ATP.

Die Protonen und die Elektronen des $\text{NADH} + \text{H}^+$ und des FADH_2 (insgesamt 24) werden zusammen mit 6 Molekülen O_2 , die durch die Membran in die Mitochondrienmatrix transportiert werden, zu 12 H_2O oxidiert.

Bilanz der Atmungskette (bei maximaler ATP-Ausbeute durch den Malat-Aspartat-Shuttle:

$10 \text{ NADH} + 10 \text{ H}^+ + 2 \text{ FADH}_2 + 28 \text{ ADP} + 28 \text{ P}_i + 6 \text{ O}_2 \rightarrow 10 \text{ NAD}^+ + 2 \text{ FAD} + 12 \text{ H}_2\text{O} + 28 \text{ ATP}.$

Damit können in der gesamten Zellatmung bis zu 32 Moleküle ATP gebildet werden: 2 in der Glykolyse, 2 beim Citratzyklus (in Form von GTP) und 28 in der Atmungskette.

Da Prokaryoten keine Zellkompartimente besitzen, müssen sie nicht Energie für intrazelluläre Transportvorgänge ausgeben und können aus einer Glucose 36 bis 38 ATP gewinnen.

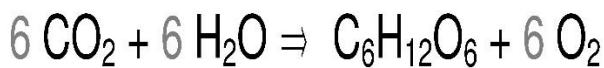
Kapitel 7: Photosynthese

Wir befassten uns im letzten Beitrag mit der Zellatmung. Bei der Zellatmung wird aus Zucker mit Hilfe von Sauerstoff Energie in Form von ATP gespeichert. Als Produkte entstehen dabei CO₂ und Wasser.

In diesem Beitrag befassen wir uns mit der Photosynthese. Diese ist quasi gesehen die umgekehrte Richtung der Zellatmung: Aus CO₂ und Wasser werden Zucker und Sauerstoff.

Die entsprechende Summenformel der Photosynthese lautet: $12 \text{ H}_2\text{O} + 6 \text{ CO}_2 \rightarrow \text{C}_6\text{H}_{12}\text{O}_6 + 6 \text{ O}_2 + 6 \text{ H}_2\text{O}$. Eine vereinfachte, verkürzte Formel reduziert die Anzahl der Wassermoleküle, da sie sich mathematisch rauskürzen lassen: $6 \text{ H}_2\text{O} + 6 \text{ CO}_2 \rightarrow \text{C}_6\text{H}_{12}\text{O}_6 + 6 \text{ O}_2$.

Photosynthese



Kohlendioxid + Wasser $\Rightarrow$ Glukose (Zucker) + Sauerstoff

Die **endotherme** Reaktion braucht **Sonnenlicht**,
das von **Chlorophyll** (Blattgrün) absorbiert wird.

Natürlich ist die Photosynthese keine bloße Umkehrung der Zellatmung. Das kann alleine deswegen schon nicht sein, weil Zucker energiereicher ist als CO₂ und Wasser. Damit aus diesen beiden Molekülen Zucker gebildet werden kann, muss also zusätzlich Energie verwendet werden. bei der Photosynthese handelt es sich um eine endotherme Reaktion, während die Zellatmung exotherm ist.

Wir erinnern uns: exotherme Reaktionen benötigen wenig Aktivierungsenergie und die Reaktionsprodukte sind am Ende energieärmer als die Ausgangsstoffe, auch als Edukte bezeichnet. Die Aktivierungsenergie wird bei der Zellatmung mittels Enzymen überwunden und die frei gewordene Energie wird in Form von ATP oder Wärme gespeichert.

Bei endothermen Reaktionen braucht es zusätzlich Energie (meist Wärmeenergie) von außen, damit die Ausgangsstoffe miteinander eine Reaktion eingehen. Die entstandenen Reaktionsprodukte sind auf einem höheren Energielevel als die Ausgangsstoffe (Edukte).

Doch woher stammt diese zusätzliche Energie, damit aus CO₂ und Wasser Zucker gebildet werden kann? Der Name Photosynthese sagt es bereits: es ist das Licht.

Licht kann aber nicht von Tieren und damit von uns Menschen als Energiequelle zur Produktion von Nährstoffen genutzt werden – auch wenn einige wahnsinnige Esoteriker behaupten Lichtwesen zu sein oder sich von Licht zu ernähren. Wir müssen von außen organisches Material aufnehmen, um uns mit Nährstoffen zu versorgen. Wenn man das tut, spricht von heterotroph.

Organismen, die über die Photosynthese Lichtenergie in chemisch gebundene Energie umwandeln können, sind alle grünen Pflanzen, die Cyanobakterien und einige Bakteriengruppen (phototrophe Bakterien). Letztere betreiben allerdings eine sog. anoxygene Photosynthese, da sie andere Substanzen als H_2O als Elektronendonoren verwenden und daher nicht O_2 freisetzen. Wir konzentrieren uns auf die oxygene Photosynthese, also jene, die Sauerstoff freisetzen.

Sind Lebewesen in der Lage aus anorganischen Substanzen organische Verbindungen aufzubauen, spricht man von autotroph. Dazu gehören – mit wenigen Ausnahmen – die Pflanzen. Aber warum sind Pflanzen dazu in der Lage?

Hierfür müssen wir uns – bevor wir uns mit der Photosynthese als solches auseinandersetzen – mit den Grundlagen befassen.

Chlorophyll und Chloroplasten

Warum sind Pflanzen eigentlich so grün? Sie verfügen über einen Farbstoff namens Chlorophyll, der ihnen die grüne Farbe verleiht (Chlorophyll heißt auf Deutsch Blattgrün). Pflanzen, Algen und Cyanobakterien besitzen verschiedene Chlorophylltypen. Sie werden als Chlorophyll a, b, c1, c2, d und f bezeichnet. Bei Bakterien, die die anoxygene Photosynthese betreiben, kommen Bakteriochlorophylle vor.

Alle Chlorophylle haben das gleiche Grundgerüst, das aus vier Pyrrolringen, dem sogenannten Porphyringerüst mit mehreren Doppelbindungen und einem im Zentrum komplex gebundenen zweifach positiven Magnesium-Ion Mg^{2+} besteht (Abb. 56). Porphyrine sind organisch-chemische Farbstoffe mit ringförmigen Kohlenstoffgerüsten. Die in Abb. 2 mit R1-4 abgekürzten Bereiche sind die Reste der verschiedenen Chlorophylle. Außerdem haben die Chlorophylle außer c1 und c2 weisen zusätzlich einen Phytolrest auf.

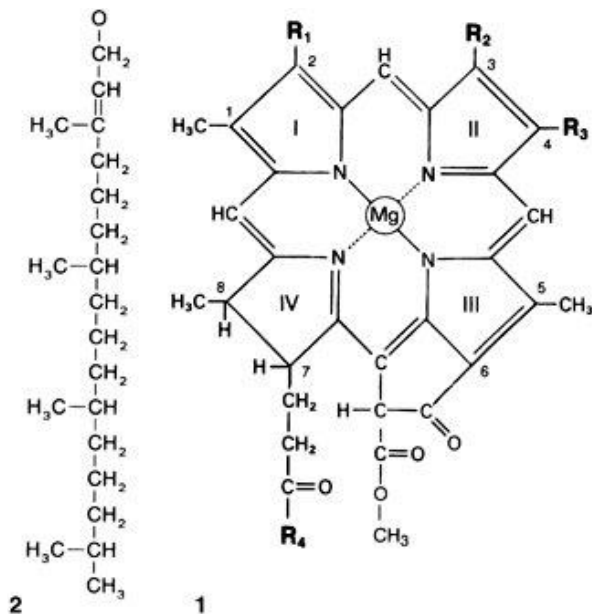


Abbildung 56 1 = Grundgerüst des Chlorophylls; 2= Phytolrest

Die Struktur des Chlorophylls ermöglicht es ihnen das Licht bestimmter Wellenlängen zu absorbieren.

Licht ist ein Teil der elektromagnetischen Wellen. Beispiele für elektromagnetische Wellen sind Radiowellen, Mikrowellen, Wärmestrahlung, Licht, Röntgenstrahlung und Gammastrahlung (Aufzählung nach aufsteigender Frequenz). Im engeren Sinne sind vom gesamten elektromagnetischen Spektrum nur die Anteile als Licht definiert, die für das menschliche Auge sichtbar sind. Im weiteren Sinne werden auch elektromagnetische Wellen kürzerer Wellenlänge (Ultraviolett) und größerer Wellenlänge (Infrarot) dazu gezählt. Weißes Licht ist eine Mischung der Farben aus dem sichtbaren Spektrum das von Violett über Blau, Grün, Gelb und Rot reicht. Diese Farben haben unterschiedliche Wellenlängen zwischen etwa 400 Nanometern beim Violett und 700 Nanometern beim Rot (Abb. 57).

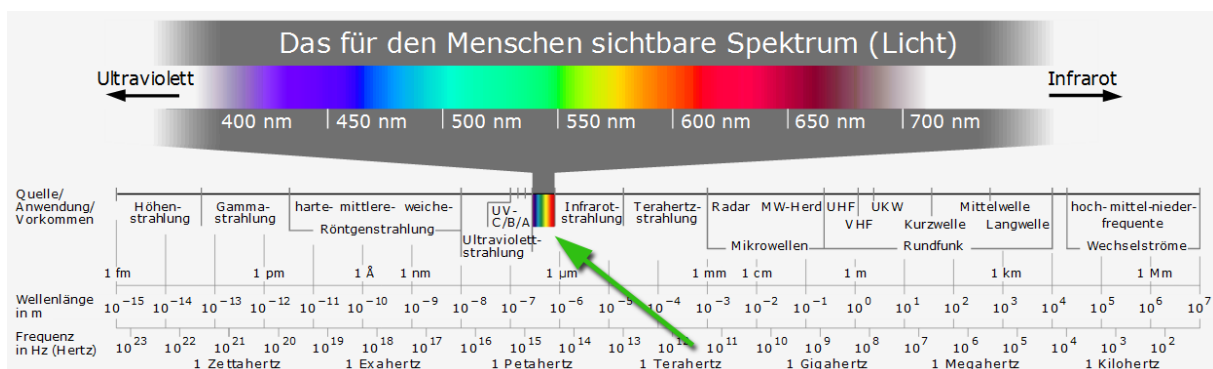


Abbildung 57 elektromagnetische Wellen.

Wenn Licht auf einen Gegenstand oder eine Substanz fällt und die Lichtstrahlen nicht reflektiert werden, dann werden sie absorbiert. Sie werden sozusagen „verschluckt“.

Seine Struktur erlaubt es, dass Chlorophyll Licht der grünen Wellenlänge reflektiert (daher erscheint es uns als grün), Licht der roten und blauen Wellenlänge absorbiert (also verschluckt).

Die verschiedenen Chlorophyllmoleküle haben Licht abgewandelte Absorptionsspektren, absorbieren Licht also etwas unterschiedlich. Abb. 58 zeigt die Absorptionsspektren von Chlorophyll a und b.

Zusammen absorbieren Chlorophyll a und b hauptsächlich im blauen Spektralbereich (400–500 nm) sowie im roten Spektralbereich (600–700 nm). Im grünen Bereich hingegen findet keine Absorption statt, so dass grünes Licht gestreut wird, was Blätter grün erscheinen lässt.

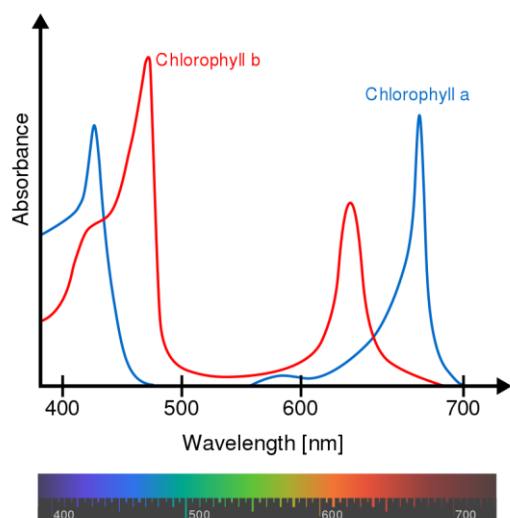


Abbildung 58 Absorptionsspektrum von Chlorophyll a und b.

Chlorophylle spielen bei der Photosynthese eine wichtige Rolle, denn sie sind in der Lage die absorbierte Energie des Lichts weiterzuleiten, also quasi als photochemischer Katalysator, welches Elektronentransportketten initiiert.

Damit solche Elektronentransportketten, wir haben ja schon welche bei der Zellatmung kennengelernt, effizient funktionieren, muss das Chlorophyll in Lichtsammelkomplexen organisiert sein. Chlorophyll „schwimmt“ nicht einfach so frei in der Pflanzenzelle herum, sondern befindet sich konzentriert in einem für sie typischen Organell: den Chloroplasten (Abb. 59).

Chloroplast

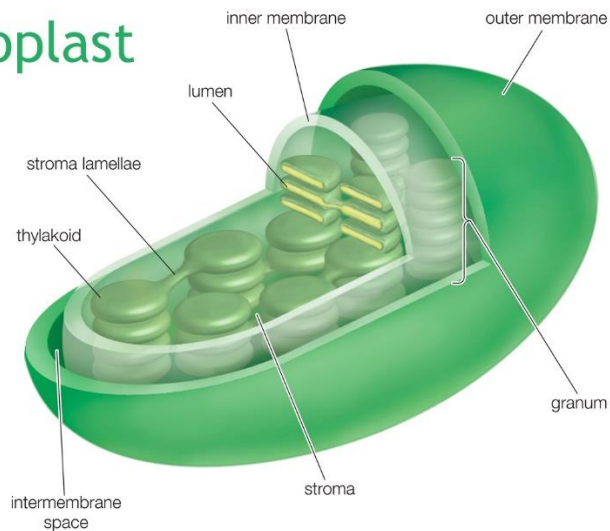
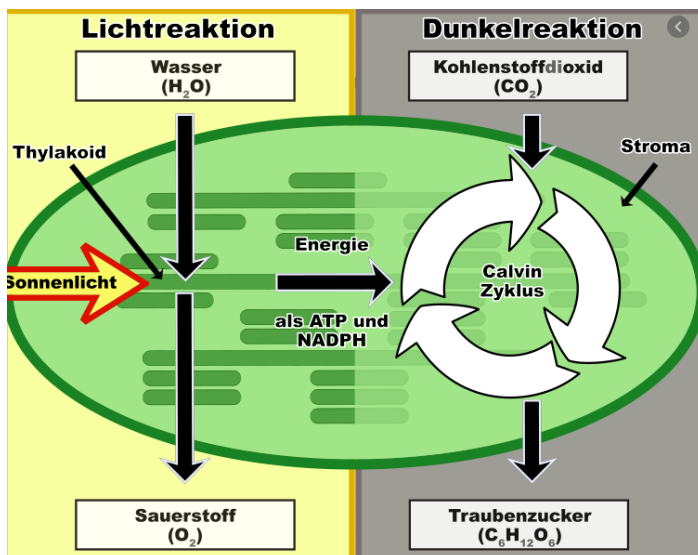


Abbildung 59: Chloroplast

Der strukturelle Aufbau der Chloroplasten gleicht dem der Cyanobakterien. Chloroplasten besitzen (fast immer) eine eigene DNA (Chloroplasten-DNA, abgekürzt cpDNA oder ctDNA) zusammen mit eigenen Ribosomen. Chloroplasten haben eine Doppelmembran. Innerhalb der inneren Membran befindet sich die Chloroplastenmatrix, auch Stroma genannt. Eingebettet in diese Matrix liegt ein internes Membransystem (die Thylakoidmembran). Dieses Thylakoid-Membransystem kann geldrollenartig übereinandergestapelt sein, was man als Granum bezeichnet. In den Membranen der Thylakoide sind verschiedene Pigmente eingelagert, vor allem der grüne Farbstoff Chlorophyll. Besonders viel davon findet sich in den Membranen der Grana, weshalb diese intensiv grün gefärbt erscheinen. Hier findet die Photosynthese statt.



Die Photosynthese kann in die beiden Teilprozesse Lichtreaktion und Dunkelreaktion untergliedert werden. Die Lichtreaktion findet in der Thylakoidmembran, die Dunkelreaktion im Stroma statt. Im Rahmen der Lichtreaktion wird die Energie bereitgestellt.

In der Dunkelreaktion wird diese Energie dazu eingesetzt, um aus CO_2 energiereiche organische Verbindungen zu synthetisieren.

Wir werden uns beide Reaktionen: Lichtreaktion und Dunkelreaktion näher betrachten.

Lichtreaktion

Die Lichtreaktion läuft in der Thylakoidmembran der Chloroplasten ab. Diese enthält zwei große Proteinkomplexe die man als Photosysteme 1 und 2 bezeichnet (Abb. 60). Sie sind mit der Fähigkeit ausgestattet, Sonnenergie einzufangen und umzuwandeln. Weiterhin befinden sich zwei Proteinkomplexe, die denen der Atmungskette der Mitochondrien ähneln: ein Cytochrom-b6-f-Komplex und die Chloroplasten-ATP-Synthase.

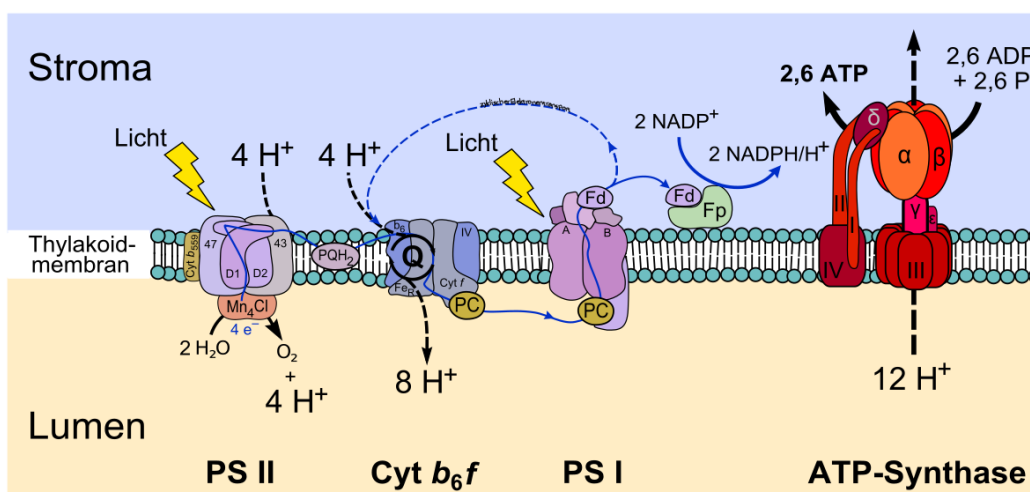


Abbildung 60 Feinbau der Thylakoidmembran

Ein Photosystem setzt sich aus einem sogenannten Antennenkomplex und aus einem Reaktionszentrum zusammen. Der Antennenkomplex besteht aus ca. 30 Proteinen, die mit Pigmentmolekülen verbunden sind. Sie werden durch das Licht in einen energiereichen, angeregten Zustand angehoben. Eine Photosynthese durch einfache Pigmente wäre relativ ineffizient, da diese dem Licht nur eine geringe Fläche entgegenstellen würden und zudem nur in einem engen Wellenlängenbereich absorbieren würden. Durch die Anordnung von chlorophyllhaltigen Lichtsammelkomplexen zu Antennen um ein gemeinsames Reaktionszentrum wird sowohl der Querschnitt vergrößert, als auch das Absorptionsspektrum verbreitert. Die eng benachbarten Chromophore in den Antennen geben die Lichtenergie von einem Pigment zum anderen weiter. Durch Exzitonen transfer kann diese Energie an das Reaktionszentrum weitergeleitet werden. Das Reaktionszentrum der Photosysteme enthält zwei Chlorophylle, die als primärer Elektronendonator fungieren. Durch die Lichtenergie wird eine Elektronentransportkette in Gang gesetzt.

Beide Photosysteme unterscheiden sich in ihren Details und Absorptionsspektren.

Das Photosystem II enthält insgesamt ca. 250 Moleküle Chlorophyll a und b sowie ca. 110 Carotinoide. Das Reaktionszentrum des Photosystems II hat ein Absorptionsmaximum bei 680 nm („P680“, Abb. 62).



Vereinfacht ausgedrückt: Sie können sich nicht entscheiden an welcher Stelle der Bindung sie bleiben sollen und wandern hin und her, sodass sich in einem Kohlenstoffring mit konjugierten, also abwechselnden, Doppelbindungen die Position

der Einfach und Doppelbindungen ändert. Man spricht auch von mesomeren Grenzstrukturen. Dies ist am Beispiel des Benzols in Abb. 63 dargestellt.

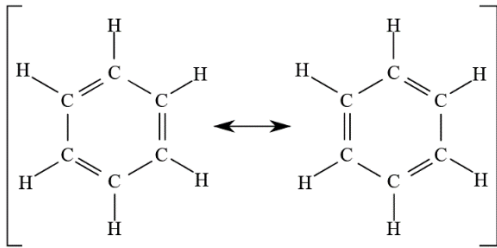


Abbildung 63 mesomere Grenzstrukturen beim Benzol.

Solche mesomeren Grenzstrukturen sind beim Chlorophyll ebenso vorhanden und das hat entsprechende Konsequenzen. Absorbiert ein Chlorophyllmolekül ein Lichtquant, führt die Energie des Lichtquants dazu, dass eines der delokalisierten Elektronen angeregt wird und von einer energieärmeren Schale bzw. Orbital (wir erinnern uns an das Schalenmodell der Atome) in ein energetisch höheres angehoben wird. Dieses angeregte Elektron im Chlorophyllmolekül strebt aber danach auf seinen Grundzustand zurückzukehren. Welche Möglichkeiten hat nun dieses angeregte „ADHS-Elektron“ seinen Grundzustand zu erreichen?

Folgende drei Möglichkeiten gibt es:

1. Umwandlung der zusätzlichen Energie in Wärme
2. Direkte Übertragung der Energie – aber nicht des Elektrons – auf ein anderes Chlorophyllmolekül in der Nachbarschaft
3. Übertragung des angeregten Elektrons auf einen Elektronenakzeptor, dann Rückkehr des nun positiv geladenen Chlorophyllmoleküls in seinen Grundzustand durch Aufnahme eines anderen, nicht angeregten Elektrons.

Die letzten beiden Fälle treten auf, wenn Chlorophyll an Proteinkomplexe gebunden ist. Das trifft auf die Photosysteme zu. Photosysteme II und I sind in Serie geschaltet und durch eine Elektronentransportkette verbunden. Werden die Redoxpotentiale aller an der Reaktion beteiligten Redoxpartner aufgetragen, ergibt sich eine Art Zick-Zack-Verlauf, der an ein gedrehtes „Z“ erinnert (Abb. 64).

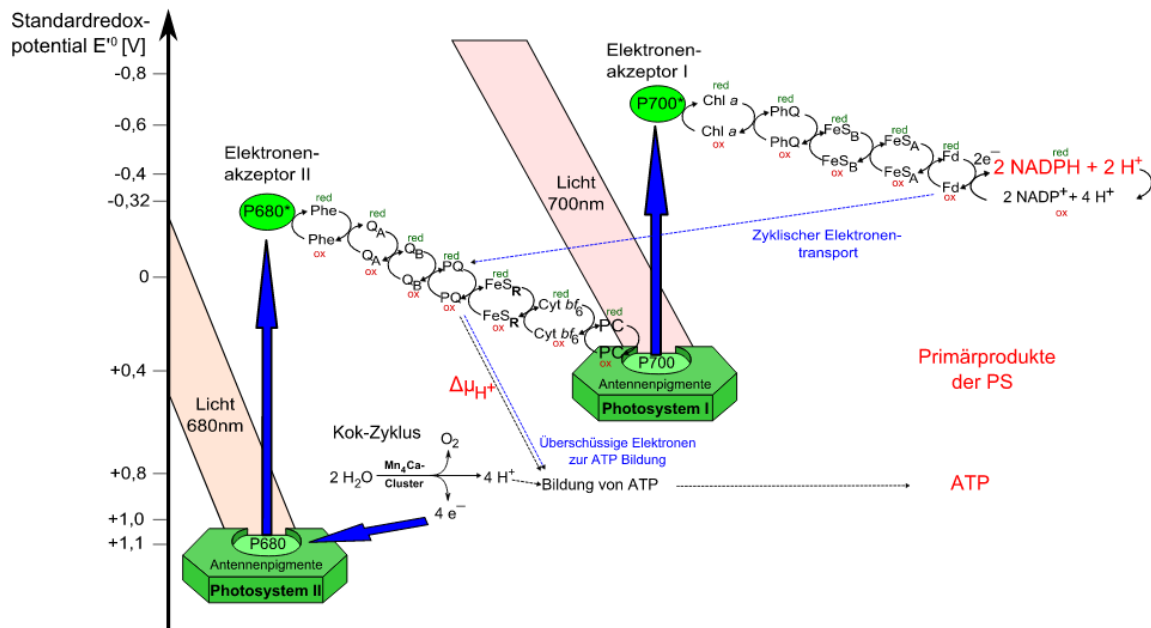


Abbildung 64 Ablauf der Lichtreaktion

Beginnen wir beim Photosystem II (diesen Namen hat er übrigens, weil er als zweites entdeckt wurde!). Die Farbstoffpigmente der Lichtsammelkomplexe transportieren die Lichtenergie weiter, bis sie das Reaktionszentrum erreicht haben.

In diesem Reaktionszentrum absorbiert das zentrale Chlorophyll-Paar schließlich so viel Energie, dass es das angeregte Elektron tatsächlich auf einen Elektronenakzeptor abgibt.

Das positiv geladene Chlorophyll-Molekül ist nun ein sehr starker Elektronenakzeptor (Oxidationsmittel), da er wieder zur Ladungsneutralität gelangen möchte. Deshalb entzieht er einem Wassermolekül das fehlende Elektron. Wasser wird aus diesem Grund durch einen wasserspaltenden Proteinkomplex in Sauerstoff, zwei Elektronen und zwei Wasserstoffprotonen (H^+) gespalten.

Hier fällt eines auf: der produzierte Sauerstoff stammt nicht aus dem CO_2 – der bei der Lichtreaktion noch keine Rolle spielt, sondern aus dem H_2O , also dem Wasser!

Das abgegebene Elektron wird nun über eine Kette an Redoxsystemen (primärer Elektronenakzeptor => Plastochinon => Cytochrom- b_6 -f-Komplex => Plastocyanin) an das Photosystem I weitergeleitet. Jedes Redoxsystem besteht aus einem Elektronenakzeptor, der ein Elektron aufnimmt und dadurch zum Elektronendonator wird. Der Elektronendonator gibt dieses Elektron dann an ein weiteres Redoxsystem mit geringerem Energiegehalt ab. Bei diesem Elektronentransport wird Energie frei, wodurch die Redoxsysteme Protonen in den Thylakoidinnenraum „pumpen“. Dazu kommen wir aber später noch einmal genauer.

Weiter geht es am Photosystem I: Das Elektron, das vom Photosystem II über die Elektronentransportkette weitergeleitet wurde, erreicht nun das Photosystem I. Es

erfolgt auch hier analog zum Photosystem II eine Anregung des sichtbaren Lichts und eine Weitergabe der Energie bis zum Reaktionszentrum.

Dort besitzt das spezielle Chlorophyll a Molekül ein Absorptionsmaximum von 700 nm. Das spezielle Chlorophyll Molekül gibt sein Elektron nun genauso wie beim PSII durch Lichtanregung ab und überträgt es auf ein weiteres Redoxsystem. Dadurch entsteht wieder eine Elektronenlücke am Chlorophyll-Molekül, die über das in der Transportkette vermittelte Elektron geschlossen werden kann.

Das letzte „Glied“ der Elektronentransportkette ist ein Enzym – die NADP⁺-Reduktase. Das kann NADP⁺ nun durch Aufnahme von 2 Elektronen und einem Wasserstoffproton (H⁺) zu NADPH reduzieren. NADP⁺ steht für Nicotin-Amid-Adenin-Dinucleotid-Phosphat, ein Coenzym, das mit dem NAD⁺ der Zellatmung verwandt ist.

Bildung von ATP

Wie wir bereits aus der Atmungskette der Zellatmung wissen, veranlasst die Elektronentransportkette die Bildung des Energieträgers ATP. Dafür ist ein Proteinkomplex zuständig – die ATP-Synthase. Dieser ist am Ende der Elektronentransportkette hinter dem Photosystem I in der Thylakoidmembran lokalisiert.

Beim Elektronentransport wird Energie frei, die die beteiligten Redoxsysteme zum Transport von Wasserstoffionen (H⁺) aus dem Stroma in den Thylakoidinnenraum nutzen. Das hat zur Folge, dass einerseits die H⁺-Konzentration im Stroma abnimmt und dadurch der pH-Wert ansteigt. Andererseits kommt es zu einer hohen H⁺-Ansammlung im Innenraum, wodurch auch der pH-Wert sinkt. Dadurch entsteht ein Konzentrationsunterschied (chemisches Potential) – in diesem Fall ein Protonengradient, da wir es hier mit Protonen zu tun haben. Zusätzlich entsteht ein Spannungsunterschied (elektrisches Potential), da die Stromaseite nun negativ und der Innenraum positiv geladen sind.

Biomembranen wie die Thylakoidmembran bestehen aber aus einer Doppelschicht aus Phospholipiden und bilden dadurch eine Barriere für die geladenen Wasserstoffprotonen. Sie sind also im Thylakoidinnenraum „gefangen“. Sie können nur durch ein Kanalprotein – die ATP-Synthase – zurück ins Stroma diffundieren, um den Konzentrations- und Ladungsunterschied auszugleichen. Daran ist die Synthese von ATP – der universellen „Energiewährung“ unserer Zellen – gekoppelt.

Der Protonenrückstrom erzeugt Energie. Das kannst du dir wie bei einem Wasserkraftwerk vorstellen, bei dem Wasser hinter einer Staumauer angestaut wird. Dieser Rückstau wird verwendet, um die Turbinen zur Rotation zu bringen, um Strom zu generieren. Diese Turbine ist unserem Fall die ATP-Synthase. Sie koppelt die Diffusion der Protonen, wie ihr Name vermuten lässt, mit der Synthese von ATP aus ADP und einer Phosphatgruppe. Da in unserem Fall Lichtenergie zu dieser Phosphorylierung führt, kannst du den Prozess auch als Photophosphorylierung bezeichnen.

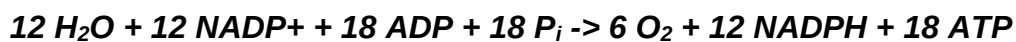
Zyklischer Elektronentransport

Neben dem gerade kennengelernten linearen Elektronentransport, auch als nichtzyklischer Elektronentransport bezeichnet, läuft unter bestimmten Bedingungen auch ein zyklischer Elektronentransport ab. Dieser besitzt, wie der Name vorgibt, einen kreisförmigen Reaktionsweg. Wichtig ist, dass du dir merkst, dass dieser Elektronentransport ausschließlich zur ATP-Bildung dient. Es entsteht weder Sauerstoff noch NADPH. Er dient also nur zur Energiegewinnung.

Einer der Gründe für diesen Transportweg, der auf den ersten Blick eher ineffizient erscheint, ist, dass oft mehr ATP benötigt wird, als über den nicht zyklischen Reaktionsweg bereitgestellt werden kann. Der zyklische Elektronentransport beginnt und endet im Photosystem I. Das durch Lichtanregung freigesetzte Elektron wird über Redoxsysteme in der Elektronentransportkette zwischen PSII und PSI übertragen und wieder zurück zum Reaktionszentrum des Photosystems I geleitet. Das darin enthaltene Chlorophyll-Molekül kann somit das Elektron aufnehmen (=Reduktion) und kehrt in seine ursprünglich ungeladene Form zurück.

Dabei ist auch ein kleines organisches Molekül beteiligt – das Plastochinon. Es veranlasst den Transport von zwei Wasserstoffprotonen in das Thylakoidlumen. Der daraus entstehende Protonengradient kann dann für die ATP-Synthese verwendet werden.

Die Bruttogleichung der Lichtreaktion kann man folgendermaßen formulieren:



Calvin Zyklus

Der Calvinzyklus ist ein Teilprozess der Photosynthese, auch Dunkelreaktion genannt. Dunkelreaktion deswegen, weil er auch ohne Licht ablaufen kann. Allerdings ist das etwas irreführend, denn ohne die vorherige Lichtreaktion kann auch keine Dunkelreaktion stattfinden. Die Lichtenergie muss dort nämlich zunächst über einen Umweg in chemische Energie – ATP und NADPH – umgewandelt werden. weiterhin ist der Calvin-Zyklus stark temperaturabhängig. Das kommt daher, dass im Calvin Zyklus viele Enzyme (=Proteine) beteiligt sind, die zu hohen Temperaturen nicht standhalten können (Denaturierung).

Der Calvinzyklus ist – ähnlich dem Zitronensäurezyklus der Zellatmung, eine im Kreis ablaufende Reaktionsfolge, die von verschiedenen Enzymen katalysiert wird. Er wurde nach seinem Entdecker, dem amerikanischen Biochemiker Melvin Calvin bekannt und findet bei Pflanzen im Stroma (=dem Inneren) der Chloroplasten und bei Bakterien im Cytoplasma statt.

Der Calvinzyklus ist die Fixierung des CO_2 aus der Luft und dessen Einbau in organische Moleküle.

Der Calvin-Zyklus (Abb. 65) kann in drei Phasen untergliedert werden: Die CO_2 -Fixierung, die Reduktionsphase und die Regenerationsphase. Das Schema des Calvin Zyklus verläuft folgendermaßen: Kohlenstoffdioxid aus der Luft diffundiert in Pflanzen durch Poren – sogenannte Spaltöffnungen oder Stomata – in die Blätter und von da in die Chloroplasten. Dort findet dann der Calvin Zyklus statt. Verschiedene Enzyme sorgen nun dafür, dass die Kohlenstoffatome des Kohlenstoffdioxid in einer im Kreis ablaufenden Reaktionsfolge verbaut (=Kohlenstoffdioxidfixierung) und daraus stabile energiereiche Verbindungen wie Zuckermoleküle hergestellt werden können.

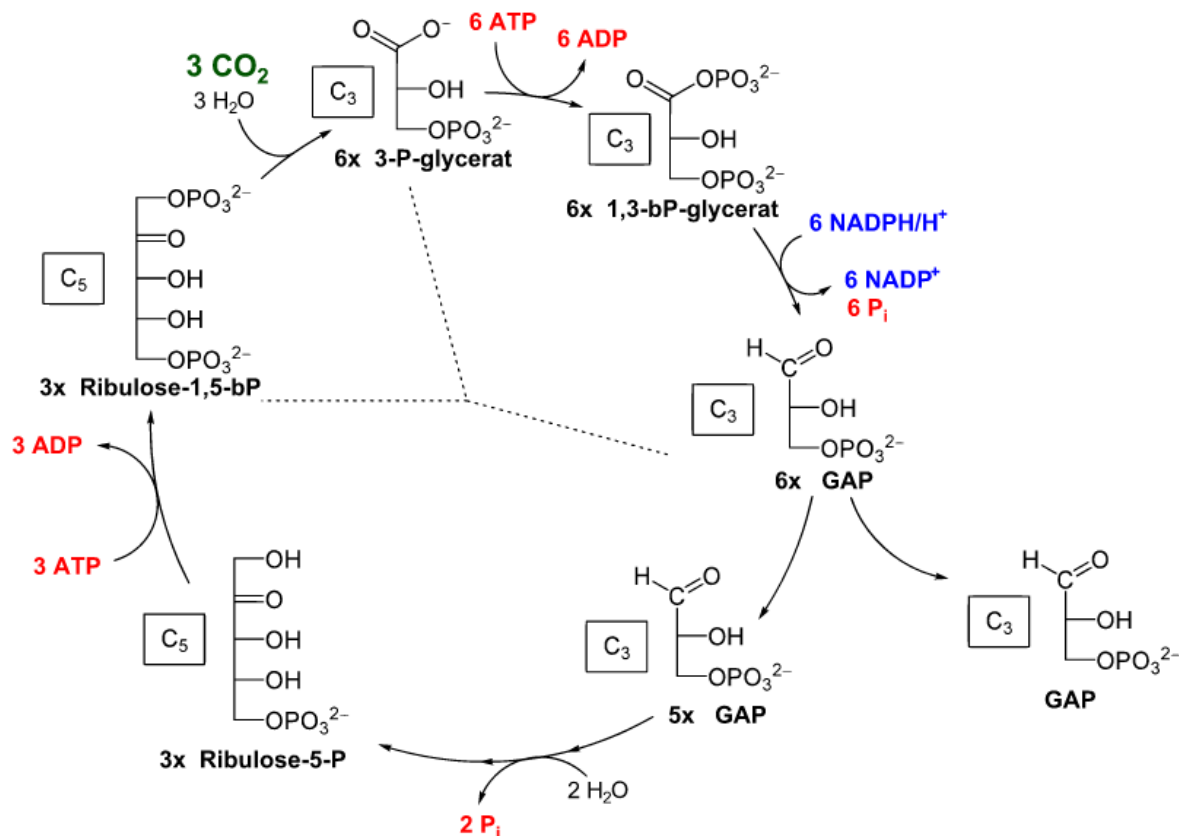


Abbildung 65 Calvin-Zyklus

CO₂- Fixierungsphase

In der Fixierungsphase verbindet sich ein Kohlenstoffdioxid-Molekül mit einem spezifischen Akzeptormolekül – dem aus 5 Kohlenstoffatomen bestehenden Zucker Ribulose-1,5-bisphosphat (RubP). Diese Reaktion katalysiert ein Enzym mit dem Namen RuBisCO, was für Ribulose-1,5-bisphosphat-Carboxylase-Oxygenase steht.

Hierbei entsteht ein C₆- Körper, der instabil ist und sofort in zwei C₃-Körper zerfällt. Sie werden auch als 3-Phosphoglycerinsäure (3-PGS) oder 3-Phosphorglycerat bezeichnet. Aus einem Kohlenstoffdioxid-Molekül entstehen also zwei 3-PGS Moleküle.

Reduktionsphase

Im nächsten Schritt – der Reduktionsphase – findet eine chemische Reduktionsreaktion statt, es werden also Elektronen aufgenommen. 3-Phosphoglycerinsäure reagiert zu einer ebenfalls aus 3 Kohlenstoffatomen bestehenden Verbindung namens Glycerinaldehyd-3-phosphat (GAP).

Da es sich um einen energieaufwendigen Prozess handelt, kann dieser nur ablaufen, wenn Energie von außen bereitgestellt wird. Da kommt jetzt die Lichtreaktion ins Spiel. Sie stellt nämlich die aus Lichtenergie gewonnenen Substrate ATP und das NADPH zur Verfügung.

Schauen wir uns nun die Reduktion im Detail an: Sie verläuft in zwei Teilschritten. Zunächst wird die 3-Phosphoglycerinsäure aktiviert, indem ein Enzym eine Phosphatgruppe auf die Säuregruppe überträgt. Die Phosphatgruppe stammt vom Energieträger ATP, der in ADP und Phosphat gespalten wird. Da die Phosphoglycerinsäure nun 2 Phosphatgruppen besitzt, wird sie jetzt als 1,3-Bisphosphoglycerinsäure bezeichnet.

Daraufhin erfolgt die eigentliche Reduktionsreaktion. Die 1,3-Bisphosphoglycerinsäure nimmt nun 2 Elektronen (e^-) und ein Wasserstoffproton (H^+) von NADPH auf. Dabei wird die zuvor übertragene Phosphatgruppe wieder abgespalten und aus der Carbonsäuregruppe ($-COOH$) entsteht nun eine Aldehydgruppe ($-CHO$). Wir erhalten eine Verbindung mit dem Namen Glycerinaldehyd-3-phosphat (GAP). Gleichzeitig wird NADPH zu $NADP^+$ oxidiert.

Sowohl NADPH als auch ATP stehen dem Calvin Zyklus durch die vorherige Lichtreaktion zur Verfügung. Deshalb kann die Dunkelreaktion auch ohne Licht nicht lange ablaufen. Gleichzeitig stehen der Lichtreaktion nach ablaufender Dunkelreaktion nun wieder ausreichend ADP und $NADP^+$ zur Verfügung. Diese beiden Reaktionen sind also voneinander abhängig.

Regenerationsphase

Der letzte Reaktionsschritt – die Regenerationsphase – dient dazu, das Akzeptormolekül (Ribulose-1,5-bisphosphat) wiederherzustellen, also zu regenerieren. Dafür werden etwa 5/6 der gebildeten C3-Moleküle Glycerinaldehyd-3-phosphat (G3P) verwendet. Auch hier ist wieder Energie in Form von ATP-Molekülen nötig. Jetzt kann der Zyklus wieder erneut durchlaufen werden, da das Akzeptormolekül wieder zur Verfügung steht.

1/6 des gebildeten Glycerinaldehyd-3-phosphat (G3P) aber verlässt den Zyklus und kann gemeinsam mit einem weiteren Molekül Glycerinaldehyd-3-phosphat (G3P) zu einem aus 6 Kohlenstoffatomen bestehenden Körper reagieren. Nach einigen Umwandlungsreaktionen entstehen daraus die Zucker Glucose (Traubenzucker) oder Fructose (Fruchtzucker). Sie können dann in den Stoffwechsel eingeschleust und dort abgebaut werden. Dadurch gewinnt die Pflanze Energie, die sie zum Beispiel für das Pflanzenwachstum benötigt.

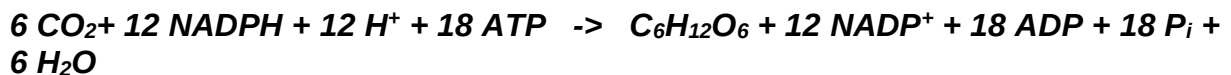
Unter günstigen Bedingungen, also wenn die Pflanzen mehr Einfachzuckermoleküle herstellt, als sie für ihren eigenen Energiebedarf benötigt, können die Zuckermoleküle als Baustoffe für verschiedene Makromoleküle wie Kohlenhydrate, Fette oder Proteine zur Verfügung. Die Kohlenhydrate können im Stroma der Chloroplasten in Stärkekörnern gespeichert und bei Bedarf mobilisiert werden.

Calvin Zyklus Bilanz

Schauen wir uns nun die Bilanz des Calvin Zyklus an: Ein Glucose-Molekül ($C_6H_{12}O_6$) besteht aus 6 Kohlenstoffatomen, wofür also 6 Kohlenstoffdioxid-Moleküle benötigt werden. Dafür sind 6 Umdrehungen des Calvin-Zyklus notwendig und ein Energieaufwand von 18 ATP Molekülen und 12 NADPH-Molekülen.

Die Bruttogleichung der Dunkelreaktion kannst du deshalb folgendermaßen formulieren:

Bruttogleichung der Dunkelreaktion



CO₂, Photosynthese, Pflanzenwachstum und Klimawandel

Wir haben uns intensiv mit den biochemischen Grundlagen der Photosynthese befasst. Dort lernten wir, dass Pflanzen aus CO₂ und Wasser Zucker und Sauerstoff (O₂) herstellen können.

Nun wird von einigen "Klima-Querdenkern" gerne folgende Behauptung aufgestellt:

„CO₂ ist kein Schadstoff. Es ist wesentlich für alles Leben auf der Erde. Die Photosynthese ist ein Segen. Mehr CO₂ schont die Natur und die Erde: Zusätzliches CO₂ in der Luft hat das Wachstum der globalen Pflanzenbiomasse gefördert. Es ist auch gut für die Landwirtschaft und erhöht die Ernteerträge weltweit.“

CO₂ ist Pflanzennahrung. Das stimmt. Ist CO₂ ein Schadstoff?

Das kommt auf die Konzentration an. Für Pflanzen ist CO₂ Nahrung, würde die Atmosphäre aber nur aus CO₂ bestehen, würden wir nicht überleben können. Schadstoffe sind immer eine Frage der Konzentration. Daher ist die Frage berechtigt: Ist mehr CO₂ automatisch gut (für Pflanzen)?

Vielleicht hilft ein Vergleich: Wir Menschen müssen Nährstoffe aufnehmen in Form von Kohlenhydraten, Proteinen und Fetten. Ein wichtiger Kohlenhydrat-Lieferant ist Saccharose, unser Haushaltszucker. Er schmeckt nicht nur gut, sondern liefert auch Energie. Aber was passiert, wenn wir zu viel Zucker zunehmen? Wir leiden an Karies,

Adipositas, Herz-Kreislauf-Erkrankungen und Diabetes. Ist zu viel Zucker gut für uns? Nein. Warum muss also mehr CO₂ automatisch gut für (alle) Pflanzen sein?

Pflanzen leben nicht von CO₂ alleine. Damit Pflanzen optimal wachsen können braucht sie auch für sie entsprechende Temperatur- und Feuchtigkeitsbedingungen.

Hier kommt es natürlich auch auf die Pflanzenart an. So wie der Mensch kein Regenwurm ist, ist eine Eiche kein Kaktus. Jede Pflanzenart reagiert unterschiedlich auf das Klima. Manche kommen mit Trockenheit besser klar als andere. Das hängt von vielen Faktoren ab. Die globale Erwärmung sorgt dafür, dass gerade die Temperatur- und Niederschlagsbedingungen sich so ungünstig auf das Wachstum vieler Kulturpflanzen auswirken, dass durch die erhöhte CO₂-Konzentration nicht kompensiert werden kann (vgl. **Lobell et al. 2011, Zhao & Running 2010**).

Die meisten wissenschaftlichen Studien zur CO₂-Steigerung wurden bisher nur in solchen geschlossenen Gewächshäusern oder in einzelnen Wachstumskammern durchgeführt. Erst vor kurzem haben Forscher begonnen ihre Aufmerksamkeit auf Experimente im Freien zu lenken. Diese Free-Air CO₂ Enrichment- oder "FACE"-Studien untersuchen die Pflanzen im Freiland, wenn sie erhöhten CO₂-Konzentrationen ausgesetzt sind. Die Ergebnisse dieser Studien sind bei weitem nicht so vielversprechend wie die von Gewächshausstudien. Die endgültigen Ertragswerte liegen in den Freiluftstudien im Vergleich zu Gewächshausstudien durchschnittlich um 50% niedriger (**Leaky et al. 2009, Long et al. 2006, Ainsworth 2005**). Gründe für das bessere Wachstum in Gewächshäusern liegen z. B. darin, dass Pflanzen in Gewächshäusern isolierter voneinander sind, das Wurzelwachstum regelmäßiger ist und Schädlinge nicht so leicht eindringen können.

Abgesehen hiervon spielen noch weitere Faktoren eine Rolle: Verschiedene Pflanzenarten reagieren auf erhöhte CO₂-Konzentrationen auf unterschiedliche Weise. Sie sind abhängig von der Pflanzenart, aber auch Alter, genetischer Variation, Photosyntheseleistung, Jahreszeit, atmosphärische Zusammensetzung, konkurrierende Pflanzen, Krankheits- und Schädlingsbefall, Feuchtigkeitsgehalt, Nährstoffverfügbarkeit, Temperatur und Verfügbarkeit von Sonnenlicht. Die fortgesetzte Zunahme von CO₂ wird ein starkes Treibmittel für eine Vielzahl von Veränderungen darstellen, die für den Erfolg vieler Pflanzen entscheidend sind, die natürlichen Ökosysteme beeinflussen und große Auswirkungen auf die globale Lebensmittelproduktion haben. Der weltweite Anstieg von CO₂ ist somit ein großes biologisches Experiment mit unzähligen Komplikationen, die es sehr schwierig machen, den Nettoeffekt dieses Anstiegs mit einem nennenswerten Detailgrad vorherzusagen.

Im Zusammenhang mit der Photosynthese spielt aber ein weiterer Faktor eine Rolle, der auch zeigt, dass Photosynthese nicht gleich Photosynthese ist: Die Photorespiration oder auch Lichtatmung genannt.

Photorespiration

Mit der Lichtreaktion kann die Pflanze das Sonnenlicht sammeln und die Lichtenergie in chemische Energie, gebunden als ATP, speichern. Diese wird bei der Dunkelreaktion, also dem Calvin-Zyklus, genutzt, um Kohlenhydrate herzustellen.

Beim Calvin-Zyklus spielt das Enzym Ribulose-1,5-bisphosphat-Carboxylase-Oxygenase, kurz RuBisCo, eine Schlüsselrolle (Abb. 66). Es fixiert das CO_2 aus der Luft auf den C5-Zucker Ribulose-1,5-bisphosphat (RubP). Durch weitere Umwandlungsschritte kann daraus Glucose oder andere C6-Zucker hergestellt werden.

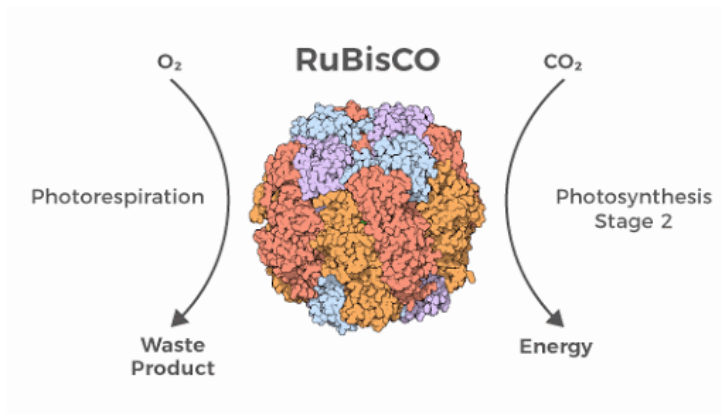


Abbildung 66 Ribulose-1,5-bisphosphat-Carboxylase-Oxygenase, kurz RuBisCo

RuBisCo hat allerdings auch noch eine weitere Eigenschaft, die man sich auch durch ihren Namen herleiten kann. Neben der Kohlenstoffdioxidfixierung (=Carboxylierung) kann es auch Sauerstoff fixieren (=Oxygenierung). Dadurch geht dem Calvin-Zyklus dann quasi ein Kohlenstoffatom „verloren“. Fixiert RuBisCo statt CO_2 Sauerstoff, entsteht ein C2-Molekül, das 2-Phosphoglycolat.

Dieses muss in einem energieaufwendigen Weg in verschiedenen Zellorganellen (Mitochondrien und Peroxisomen) zu 3-Phosphoglycerinsäure umgewandelt werden. 3-Phosphoglycerinsäure ist das Zwischenprodukt des Calvin-Zyklus, das nach der Fixierung von CO_2 entsteht. Denn nur dieses Molekül kann der Calvin-Zyklus weiter verstoffwechseln. Das C2-Molekül 2-Phosphoglycolat ist für den Calvin-Zyklus unbrauchbar, man kann mit ihm also kein Zucker herstellen.

Den Weg zur Wiederverwertung des Kohlenstoffgerüsts wird auch als Photorespiration oder Lichtatmung bezeichnet (Abb. 67). Da er ziemlich viel Energie „verschwendet“, versuchen ihn Pflanzen nach Möglichkeit zu vermeiden.

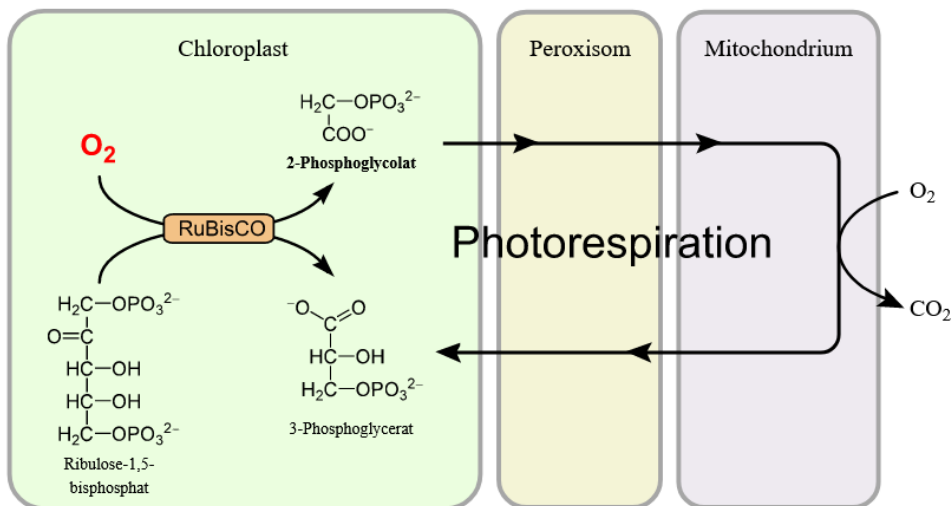


Abbildung 67 Photorespiration

Am idealsten ist, wenn RuBisCO möglichst wenig mit Sauerstoff in Berührung kommt und eher das CO_2 fixiert. Nun könnte man einwenden: Hey dann ist mehr CO_2 in der Atmosphäre doch gut für die Pflanzen. Nicht unbedingt. Mehr CO_2 in der Luft sorgt für höhere Temperaturen und weniger Luftfeuchtigkeit.

Besonders an heißen, trockenen Tagen kommt es häufig zur Lichtatmung, da die Pflanzen ihre Spaltöffnungen schließen, damit kein Wasser verdunsten kann. Dadurch kann aber auch wenig Kohlenstoffdioxid aus der Luft in die Pflanzen gelangen und der Sauerstoffgehalt steigt im Vergleich zum Kohlenstoffdioxidgehalt in der Pflanze. Das wiederum führt dazu, dass das Enzym RubisCO auch öfter Sauerstoff an das Akzeptormolekül bindet, die Lichtatmung also stattfindet.

Also bei Trockenheit geschieht es, dass die Spaltöffnungen der Pflanzen sich schließen, sich in der Pflanze der Sauerstoff ansammelt und immer weniger CO_2 aufgenommen wird und das Enzym RuBisCO nutzt dann eher den Sauerstoff zur Bildung von 2-Phosphoglycolat, also dem C2-Zucker. Da dieser für den Calvinzyklus „nutzlos“ ist, wird er unter Energieverbrauch in 3-Phosphoglycerinsäure umgewandelt und in den Calvinzyklus zur Produktion von Zucker wieder integriert. Sind die Temperaturverhältnisse ungünstig hoch bleiben die Spaltöffnungen länger geschlossen, was dann dazu führen kann, dass die Pflanzen mehr Energie verbrauchen als sie erzeugen – oder anders ausgedrückt: die Pflanze stirbt ab! Wie lange Pflanzen die Lichtatmung verwenden können ist natürlich von Art zu Art unterschiedlich.

Photosynthese Typen

Es gibt aber im Pflanzenreich verschiedene Strategien die Probleme der Lichtatmung zu umgehen. Es gibt nämlich verschiedene Photosynthese-Typen unterscheiden: die C₃- Pflanzen, die C₄- Pflanzen und die CAM-Pflanzen (Abb. 68).

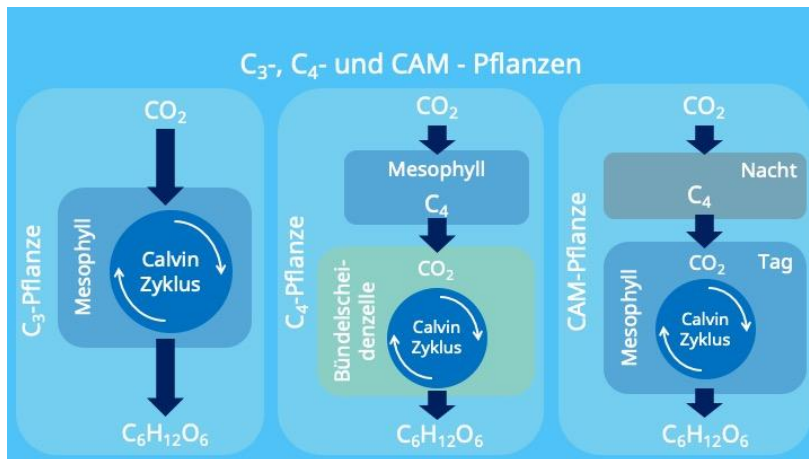


Abbildung 68 C₃, C₄ und CAM Pflanzen

C₃ Pflanzen

C₃- Pflanzen sind die ganz „normalen Pflanzen“ ohne besondere Anpassungen hinsichtlich der Vermeidung zur Photorespiration verstehen. Das bedeutet, dass der Calvin-Zyklus, so wie wir ihn im [Kapitel zur Photosynthese beschrieben](#) haben, abläuft.

Der Name C₃ kommt daher, da das erste Produkt nach der Kohlenstofffixierung, also der Anlagerung von Kohlenstoffdioxid an das Akzeptormolekül (1,5-Ribulosebisdiphosphat) im Calvin Zyklus, ein C₃ Körper (3-PGS) ist.

Die meisten Pflanzen auf der Erde betreiben diese Art der Photosynthese, darunter alle Bäume und auch Soja- oder Reispflanzen.

C₄ Pflanzen

Die C₄- Pflanzen haben jetzt eine Strategie entwickelt, um dem verschwenderischen Weg der Lichtatmung zu entgehen. Ziel dieser Strategie ist es, dass das Enzym RuBisCO immer einer hohen Konzentration an Kohlenstoffdioxid ausgesetzt ist, damit es nicht die Möglichkeit hat, an Sauerstoff zu binden (Abb. 69 und 70).

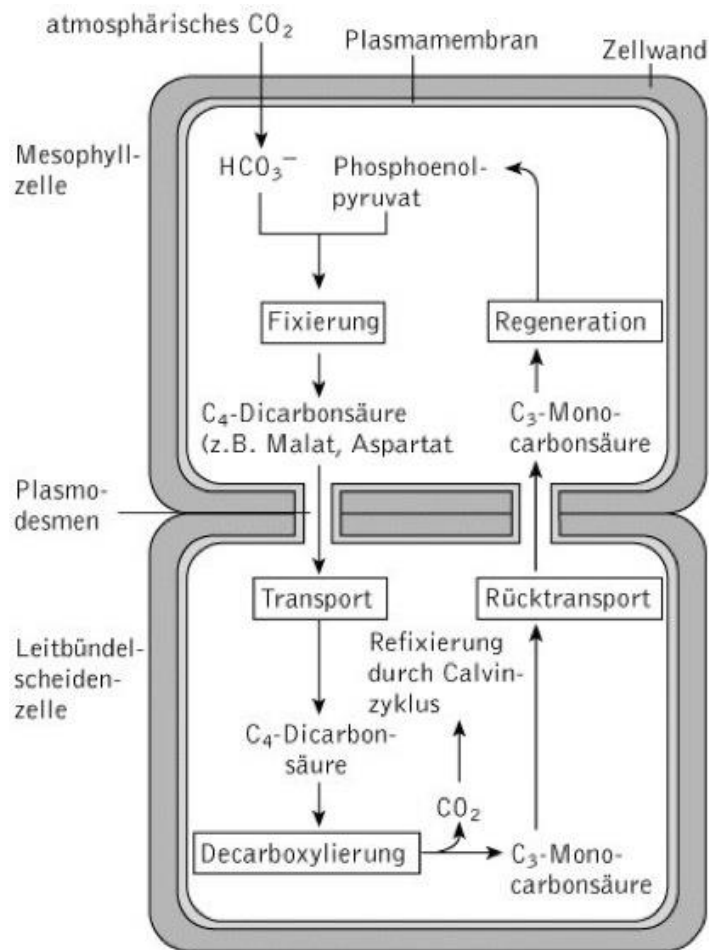


Abbildung 69 C₄-Photosynthese. 1. CO_2 wird vorfixiert. 2. Die daraus gebildete C₄-Verbindung wird in das benachbarte Kompartiment transportiert. 3. Dort wird CO_2 freigesetzt und im Calvin-Zyklus verbraucht. 4. Der übrige C₃-Körper wird wieder zurücktransportiert und für einen weiteren Zyklus regeneriert.



Abbildung 70 typische C₄-Pflanzen

Die spezielle Anpassung der C₄-Pflanzen ist eine räumliche Trennung zwischen der Kohlenstoffdioxid-Fixierung und dem eigentlichen Calvin-Zyklus.

Es findet hier nämlich eine Vorfixierung von Kohlenstoffdioxid statt. In den sogenannten Mesophyllzellen wird das CO₂ mit dem Phosphoenolpyruvat (PEP), ein energiereiches Stoffwechselzwischenprodukt der Glykolyse, fixiert, sodass als Zwischenprodukt Oxalacetat entsteht. Oxalacetat ist eine organische Säure mit 4 Kohlenstoffatomen. Daher kommt auch der Name C4– Pflanzen zustande. Oxalacetat wird dann Dicarbonsäure wie Malat oder Asparat umgewandelt.

Die entstehende Säure kann dann weiter in spezielle Zellen (Bündelscheidenzellen) transportiert werden, die sich wie eine Art Kranz anordnen. Dort findet die Kohlenstoffdioxidspaltung der Säure statt und der Calvin-Zyklus kann dann wie gewohnt ablaufen.

Etwa 3% der Arten der Bedecktsamer betreiben eine C4-Fotosynthese und machen etwa 5% der Biomasse aus. Sie fixieren aber 23% des CO₂ (**Bond et al. 2005, Kellogg 2013**). Hierzu gehören vornehmlich Gräser, die aus wärmeren Gebieten mit längeren Trockenperioden stammen, z. B. Mais, Zuckerrohr oder Hirse (Abb. 70). C4-Photosynthese kommt aber auch in anderen Pflanzenfamilien vor und ist wahrscheinlich 65-mal unabhängig voneinander in 19 Pflanzenfamilien entstanden. Bei Bäumen kommt die C4-Photosynthese mit wenigen Ausnahmen praktisch gar nicht vor (**Sage 2017, Bresinski et al. 2008, S. 312**).

Gegenüber C3-Pflanzen zeichnen sich C4-Pflanzen bei Wasserknappheit, hohen Temperaturen und Sonneneinstrahlung aus. C4-Pflanzen dominieren also eher trockene Gebiete, wie wir sie z. B. aus den Savannen Afrika oder Australien kennen (Abb. 7). So betreiben etwa 70 % aller im Death-Valley-Nationalpark lebenden Arten eine C4-Photosynthese (**Bresinski et al. 2008, S. 311**).

In tropischen Regenwäldern dominieren die C3-Pflanzen (Abb. 71), was nicht verwundern sollte, da die C4-Photosynthese in Bäumen praktisch gar nicht vorkommt. Bei Nadelbäumen, Farnen und Moosen kommt eine C4-Photosynthese gar nicht vor (**Sage 2017**).

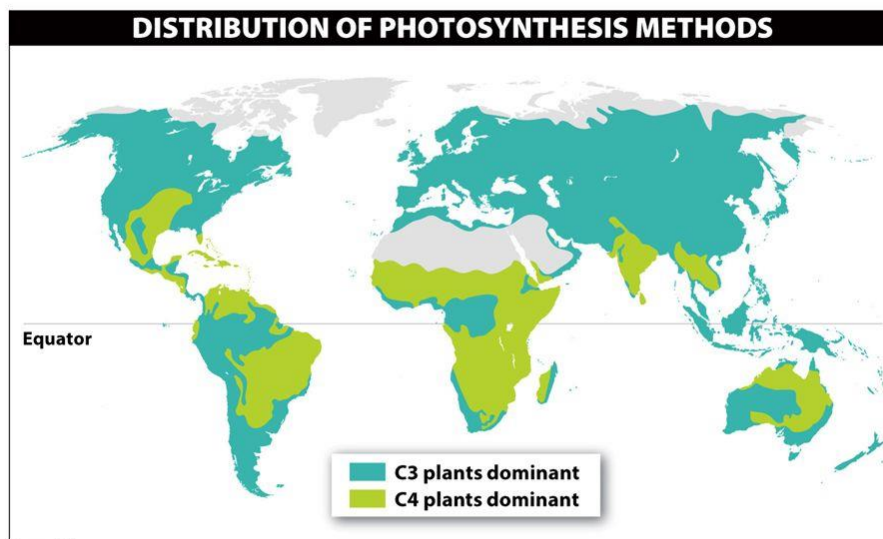


Figure 4-26
What Is Life? A Guide To Biology, Third Edition
© 2015 W. H. Freeman and Company

Abbildung 71 Verbreitung der C3 und C4-Pflanzen.

CAM Pflanzen

Auch die CAM-Pflanzen (=Crassulaceen-Säurestoffwechsel) besitzen eine spezielle Strategie, um die Lichtatmung zu minimieren. Hier erfolgt eine zeitliche Trennung zwischen der Kohlenstoffdioxidfixierung und dem Ablauf des Calvin Zyklus (Abb. 72).

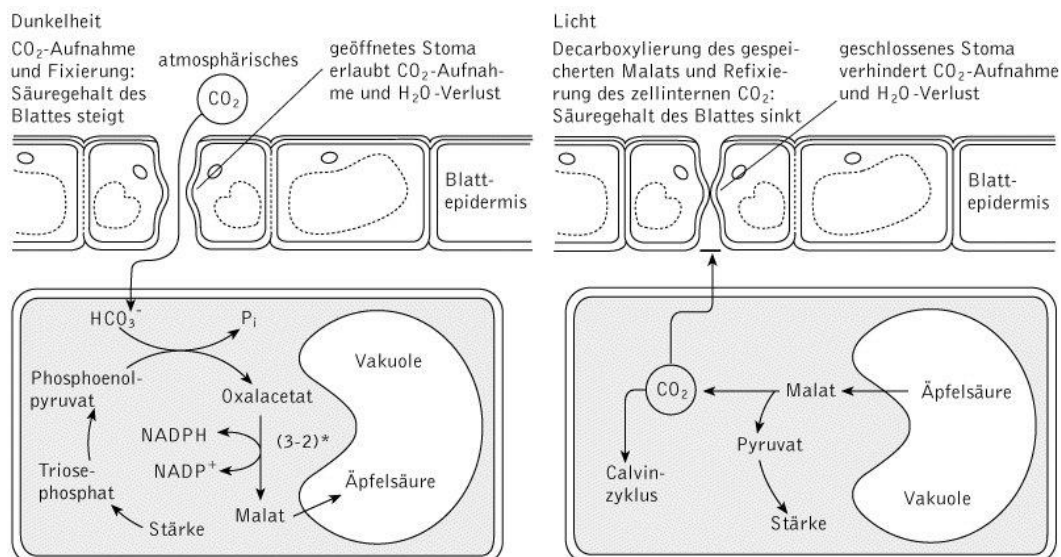


Abbildung 72 CO₂-Fixierung bei CAM-Pflanzen

Das läuft so ab: In der Nacht öffnen die Pflanzen ihre Spaltöffnungen, damit CO₂ in die Blätter diffundieren kann. Dort kann Kohlenstoffdioxid in eine organische Säure mit eingebaut werden (ähnlich zu den C4 Pflanzen). Die dadurch entstehende Säure aus

4 Kohlenstoffatomen (meistens Äpfelsäure) kann dann in die Vakuolen der Zellen transportiert und dort gespeichert werden.

Tagsüber schließen die Pflanzen wieder ihre Spaltöffnungen, um den Wasserverlust zu verringern. Die Säure kann nun aus den Vakuolen transportiert und Kohlenstoffdioxid abgespalten werden. Dadurch kann der Calvin-Zyklus wieder wie gewohnt ablaufen. Auch hier ist Ziel dieser Strategie, dass das Enzym RuBisCO immer einer hohen Konzentration an Kohlenstoffdioxid ausgesetzt ist, damit es nicht die Möglichkeit hat, an Sauerstoff zu binden.

Zu den CAM-Pflanzen gehören nicht nur die sukkulenten Dickblattgewächse (Crassulaceae), nach denen dieser Typ der CO₂-Fixierung benannt ist, sondern auch viele Arten aus den Familien Cactaceae, Agavaceae und Euphorbiaceae. Die Ananas ist ebenfalls eine CAM-Pflanze. Eine Reihe von Pflanzen sind zudem in der Lage, den CAM-Stoffwechsel fakultativ zu betreiben. So wechselt die Eispflanze, *Mesembryanthemum crystallinum* (Aizoaceae), bei Wassermangel vom C3-Weg (C3-Pflanzen) zum CAM-Weg. CAM-Pflanzen kommen vorwiegend in Trockengebieten vor, sind aber auch in tropischen Regenwäldern zu finden. Denn viele Aufsitzerpflanzen, sog. Epiphyten, also Pflanzen, die auf anderen Pflanzen wachsen, leiden auch unter Wassermangel und hohen Temperaturen. Hierzu zählen z. B. viele Orchideen und Bromelien. Entsprechend kommt auch bei ihnen die CAM-Photosynthese relativ häufig vor (Abb. 73).



Abbildung 73 typische CAM-Pflanzen

Weitere Photosynthese-Faktoren

Abhängigkeit von der CO₂-Konzentration

Der Kurvenverlauf für die Abhängigkeit der Fotosyntheseleistung von der CO₂-Konzentration entspricht einer Sättigungskurve (Abb. 74). Das heißt zu einem bestimmten Zeitpunkt wird ein Maximum der Photosynthese-Leistung erreicht und die weitere Zugabe von CO₂ verstärkt die Photosynthese-Leistung nicht.

Abhängigkeit von der CO₂-Konzentration

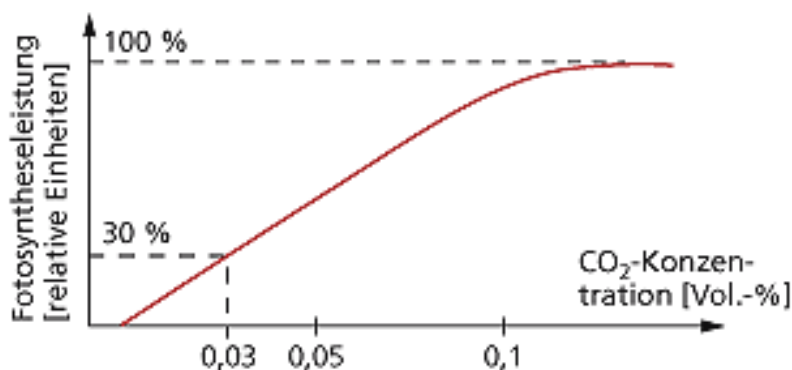


Abbildung 74 Abhängigkeit der Photosyntheseleistung von der CO₂-Konzentration.

Die Photosynthese-Leistung steigt bei den meisten Pflanzen proportional mit der Erhöhung des CO₂-anteils bis auf ca. 0,1 Vol.-%. Danach folgt bei vielen Pflanzen trotz steigender CO₂-konzentration keine Zunahme der Photosynthese-Leistung (CO₂-sättigung) mehr, wobei ein zu hoher Kohlenstoffdioxidanteil auch zu Schädigungen führen kann (hieraus erkennt man schon, dass eine ewige Steigerung des CO₂ in der Atmosphäre nicht unbedingt zum Pflanzenwachstum führen muss, denn irgendwann ist auch eine Pflanze satt).

Der CO₂-Kompensationspunkt gibt an, ab welchem Punkt die CO₂-Aufnahme und CO₂-Abgabe, also die Leistung von Photosynthese und Atmung in einer Pflanze, im Gleichgewicht sind. Bei C₃-Pflanzen liegt dieser – je nach Art – zwischen 0,005 und 0,01 Vol.-% CO₂ in der Atmosphäre.

Unter natürlichen Bedingungen ist bei C₃-Pflanzen der CO₂-Anteil in der Atmosphäre (0,03 Vol.-%) der begrenzende Faktor für die Photosynthese-Leistung. Die Photosynthese-Intensität kann bei diesen Pflanzen durch CO₂-Anreicherung in Gewächshäusern bis zum Dreifachen gesteigert werden. Wir haben aber schon vorhin kennengelernt, dass die Situation in der freien Natur eine andere ist, als in Gewächshäusern.

Photosynthese-Spezialisten wie C₄- und CAM-Pflanzen erhöhen in Anpassung an sonnige trockene Standorte unter Energieverbrauch die Konzentration an CO₂, indem sie zunächst einen anderen Akzeptor (nämlich das Phosphoenolpyruvat PEP) zur

CO₂-Fixierung verwenden. So können sie auch bei sehr niedrigen CO₂-Konzentrationen noch eine hohe Photosynthese-Leistung erzielen.

Chlorophyllgehalt

Der Chlorophyllgehalt einer Pflanze ist unter natürlichen Bedingungen kein begrenzender Faktor. So können selbst Pflanzen mit geringem Chlorophyllgehalt ausreichend Photosynthese betreiben. Der Chlorophyllgehalt kann aber bei Pflanzen in Anpassung an unterschiedliche Lichtintensitäten variieren. So weisen Schattenblätter im Vergleich zu Sonnenblättern einen höheren Chlorophyllgehalt pro Einheit Blattfläche auf (Abb. 75).

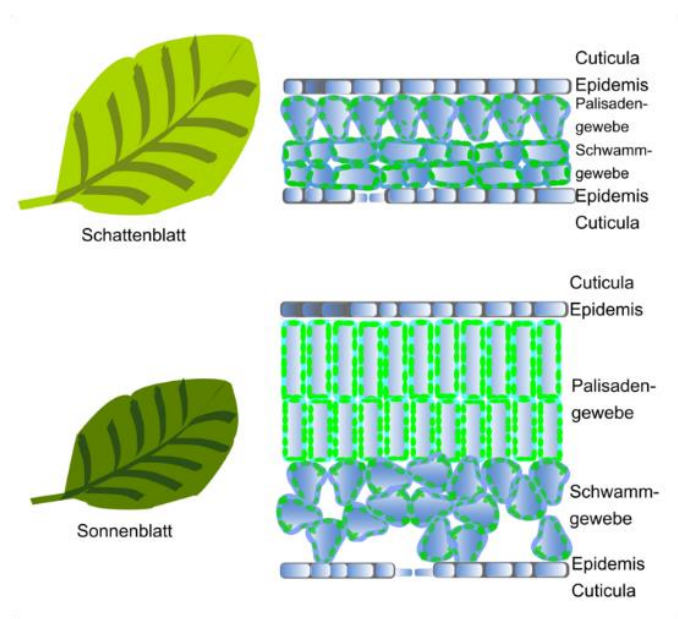


Abbildung 75 Unterschied Sonnenblatt, Schattenblatt.

Licht

Die Photosyntheseleistung hängt natürlich von der Lichtmenge und -intensität ab, weil die für den Calvinzyklus notwendigen ATP und NADPH-Moleküle in der Lichtreaktion erzeugt werden.

Abb. 76 zeigt die Abhängigkeit der Photosyntheseleistung von der Lichtintensität:

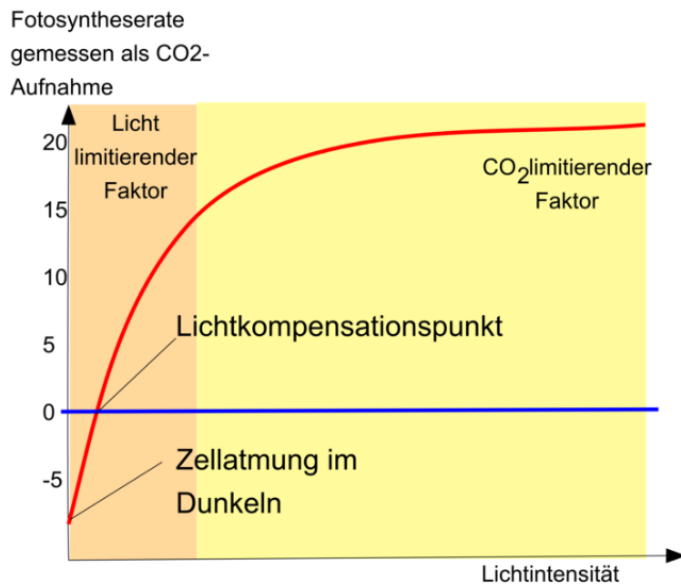


Abbildung 76 Abhängigkeit der Photosyntheseleistung von der Lichtintensität

Der Kurvenverlauf für die Abhängigkeit der Fotosyntheseleistung von der Lichtenergie beginnt im negativen Bereich, da hier vorerst der Sauerstoffverbrauch durch die Zellatmung größer ist als die Sauerstoffproduktion durch die Fotosynthese. Solange keine anderen Umweltfaktoren (z. B. Kohlenstoffdioxid, Temperatur) begrenzend wirken, wächst zunächst die Fotosyntheseintensität mit Erhöhung der Lichtintensität proportional. Am Lichtkompensationspunkt, wo die Sauerstoffabgabe (bzw. Kohlenstoffdioxidaufnahme) durch die Fotosynthese genauso so groß ist wie der Sauerstoffverbrauch (bzw. Kohlenstoffdioxidproduktion), schneidet die Kurve die x-Achse und die Nettofotosynthese hat den Wert null.

Nur Pflanzen, deren Produktion von Sauerstoff bzw. Verbrauch an Kohlenstoffdioxid über dem Lichtkompensationspunkt liegt (höhere Fotosyntheseleistung als Atmungsintensität), können organische Stoffe speichern. Mit zunehmender Strahlungsintensität wirken die anderen Faktoren nach dem Gesetz des Minimums immer stärker limitierend auf die Fotosyntheseleistung. Dadurch steigt die Kurve der Fotosyntheseleistung immer weniger und zwar bis zu dem Punkt, an dem sich trotz steigender Lichtintensität die Fotosyntheseleistung nicht mehr erhöht (Lichtsättigung).

Der Kurvenverlauf der Fotosyntheseintensität in Abhängigkeit von der Lichtintensität wird dementsprechend als Sättigungskurve bezeichnet. Falls die Lichteinwirkung trotzdem weiter ansteigt, kann dies zu einer Schädigung der Fotosysteme führen. Es werden nämlich sog. Sauerstoffradikale (O_2^-) gebildet. Als Radikale bezeichnet man Atome oder Moleküle mit mindestens einem ungepaarten Elektron, die meist besonders reaktionsfreudig sind. Sauerstoffradikale sind schädliche Formen des Sauerstoffs, die auch ein Nebenprodukt der Atmungskette der Mitochondrien sind. Werden zu viele dieser radikale gebildet, spricht man vom oxidativen Stress und sie spielen bei verschiedensten Erkrankungen sowie beim Alterungsprozess eine wesentliche pathophysiologische Rolle. Z. B. greifen sie Proteine oder die DNS an und zerstören sie.

Unter natürlichen Bedingungen treten bei Pflanzen Sauerstoffradikale auf, z. B. bei sehr geringen Temperaturen oder bei Schattenpflanzen, wenn sie plötzlich sehr

starkem Lichteinfluss ausgesetzt werden, auf. In begrenztem Umfang können bestimmte Fotosynthesepigmente (Xanthophylle) überschüssige, nicht verwertbare Lichtenergie zum Schutz vor entstehenden Sauerstoffradikalen auffangen.

Pflanzenarten zeigen unterschiedliche Anpassungen an die Lichtintensität. So unterscheiden sich Fotosynthesespezialisten, wie der Mais als C4-Pflanze, von den C3-Pflanzen. C4-Pflanzen erreichen selbst bei hoher Lichtintensität keine absolute Lichtsättigung. Deshalb sind C4-Pflanzen den C3-Pflanzen bei voller Sonneneinstrahlung in Bezug auf die Strahlungsverwertung überlegen. Bei C3-Pflanzen unterscheidet man an Starklicht (Sonne) und Schwachlicht (Schatten) angepasste Pflanzen oder Organe.

Temperatur

Wie bei allen enzymatisch gesteuerten Reaktionen ist auch die Fotosyntheseleistung von der Temperatur abhängig und wird in Form einer Optimumskurve beschrieben (Abb. 77).

Abhängigkeit von der Temperatur

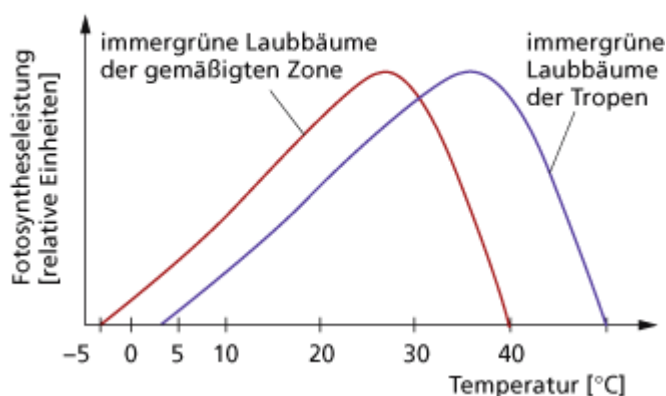


Abbildung 77 Abhängigkeit der Photosynthese von der Temperatur

So steigt die Fotosyntheseleistung bei ausreichender Versorgung mit den anderen Bedingungen entsprechend der RGT-Regel bis zum Optimum (RGT-Regel: Erhöhung der Temperatur um 10 K bewirkt eine Leistungssteigerung auf ca. das Doppelte). Nach Erreichen des Optimums wirken andere Faktoren immer stärker limitierend auf die Fotosyntheseleistung, so u. a. Licht- und Kohlenstoffdioxidversorgung, Erhöhung der Atmungsintensität oder Denaturierung der beteiligten Enzyme. Gleichzeitig haben höhere Temperaturen Einfluss auf den Wassergehalt innerhalb der Pflanze. Durch die erhöhte Transpiration sinkt der Wassergehalt in der Pflanze, sodass sich bei höheren Temperaturen die Spaltöffnungen schließen und so die Kohlenstoffdioxidzufuhr stark behindern. Die Werte der Optimumkurve können innerhalb der Pflanzenart und zwischen den Arten sehr unterschiedlich sein und können durch Temperaturanpassungen verschoben werden.

Wasser und Mineralstoffe

Wasser ist neben Kohlenstoffdioxid ein direkter Ausgangsstoff für die Fotosynthese. Die benötigte Menge an Wasser ist jedoch so gering, dass sie sich nicht direkt auf den Stoff- und Energiewechselprozess auswirkt. Vielmehr werden indirekt durch den

Wassermangel die enzymatischen Prozesse gehemmt und durch den Verschluss der Spaltöffnungen wird die Kohlenstoffdioxidaufnahme entscheidend beeinträchtigt. Die Pflanzen greifen zunächst auf intern entstandenes Kohlenstoffdioxid der Zellatmung zurück, bevor in weiteren Stufen die Kohlenstoffdioxidzufuhr ganz gedrosselt wird und dann Schäden an Fotosyntheseapparat und Membranen auftreten.

Die Mineralstoffe, die über die Wurzel in die Pflanze aufgenommen werden, sind u. a. für den Ablauf aller biochemischen Prozesse und auch zum Aufbau von Zellstrukturen notwendig. Chloroplasten benötigen Mineralstoffe (Mg, Fe, Cu, Mn) für den Bau des Fotosyntheseapparats. Diese sind außerdem Bestandteil der Pigmente und Redoxsysteme (Elektronentransportkette).

Zusammenfassend können wir sagen: Pflanzenwachstum hängt nicht alleine von CO₂ ab, sondern ist etwas komplexer. Das sollte man immer bedenken, wenn man Aussagen hört wie CO₂ sei gut für das Pflanzenwachstum.

Literatur zur Ökologie der Photosynthese

Ainsworth, E. A., Long, S. P. (2005): What have we learned from 15 years of free-air CO₂ enrichment (FACE)? A meta-analytic review of the responses of photosynthesis, canopy properties and plant production to rising CO₂. *New Phytol.*;165(2):351-71. doi: 10.1111/j.1469-8137.2004.01224.x. PMID: 15720649.

Bond, J., Woodward, F. I., Midgley, G. F. (2005): The global distribution of ecosystems in a world without fire. *New Phytologist*. 165, Nr. 2, S. 525–538. doi:10.1111/j.1469-8137.2004.01252.x. PMID 15720663.

Bresinsky, A. et al. (2008): Strasburger – Lehrbuch der Botanik. 36. Auflage. Spektrum Akademischer Verlag; Heidelberg; ISBN 978-3-8274-1455-7, S. 312.

Kellogg, E. A. (2013): C₄ photosynthesis. In: *Current Biology*. 23, Nr. 14, S. R594–R599. doi:10.1016/j.cub.2013.04.066.

Leakey, A. D. B., Ainsworth, E. A., Bernacchi, C. J., et al. (2009): Elevated CO₂ effects on plant carbon, nitrogen, and water relations: six important lessons from FACE, *Journal of Experimental Botany*, Volume 60, Issue 10, 2859–2876, <https://doi.org/10.1093/jxb/erp096>

Lobell, D., Bänziger, M., Magorokosho, C., et al. (2011): Nonlinear heat effects on African maize as evidenced by historical yield trials. *Nature Clim Change* 1, 42–45. <https://doi.org/10.1038/nclimate1043>

Sage, R. F. (2017): A portrait of the C₄ photosynthetic family on the 50th anniversary of its discovery: species number, evolutionary lineages, and Hall of Fame. *Journal of Experimental Botany*. Band 68, Nr. 2, S. e11–e28, doi:10.1093/jxb/erx005, PMID 28110278.

Zhao, M., Running, S. W. (2010): Drought-Induced Reduction in Global Terrestrial Net Primary Production from 2000 Through 2009. *Science* Vol. 329, Issue 5994, pp. 940-943 DOI: 10.1126/science.1192666

Kapitel 8: DNA, Chromosomen, Genom

Historisches

Der Tübinger Forscher Friedrich Miescher gilt als der Entdecker der Nukleinsäuren. Er isolierte in den Sechziger- und Siebzigerjahren des 19. Jahrhunderts aus tierischen und menschlichen Geweben eine Substanz, die er Nuklein nannte, da er diese Substanz aus dem Nukleus, also dem Zellkern, isoliert hatte.

Schon zu Beginn des 20. Jahrhunderts wurde postuliert, dass Nukleinsäuren Träger von Erbinformationen sind.

Doch erst 1928 konnte eine Forschergruppe um Frederick Griffith und Oswald Avery am Rockefeller Institute in New York nachweisen, dass eine chemisch relativ einfache, langkettige Nukleinsäure (Desoxyribonukleinsäure, DNS, oder in der angelsächsischen Abkürzung DNA) Träger genetischer Information sein muss.

Sie überführten eine erbliche Eigenschaft eines Bakterienstammes der Art *Streptococcus pneumoniae* auf einen anderen, ein Prozess, der heute als Transformation bezeichnet wird. *Streptococcus pneumoniae* kommt in zwei Formen vor: in einer S-Form und einer R-Form. Die S-Form ist krankheitsauslösend und die Bakterien haben eine schützende Schleimkapsel, die vom Immunsystem des Wirtes nicht erkannt wird. Der R-Form fehlt diese Schutzsubstanz und sie sind nicht krankheitserregend. Mäuse wurden mit beiden Bakterienstämmen behandelt. Erhielten die Mäuse den infektiösen S-Stamm starben die Mäuse, bei einer Injektion mit dem R-Stamm überlebten sie. Erhielten die Mäuse einen durch Hitzebehandlung abgetöteten S-Stamm überlebten sie ebenfalls. In einem weiteren Versuchsablauf wurden abgetötete S-Stämme mit lebenden R-Stämmen, die nicht infektiös sind, vermengt und den Mäusen injiziert. Das Ergebnis: die Maus starb trotzdem. Was ist geschehen? Die infektiöse Eigenschaft des S-Stammes wurde von den harmlosen R-Stamm aufgenommen, ein Prozess, der heute als Transformation bezeichnet wird. Die Substanz, die die Eigenschaften von einem Bakterienstamm auf den nächsten überträgt muss hitzebeständiger sein. 1944 verfeinerten sie ihr Experiment und konnten an den Zellen des S-Stammes unter systematischem Einsatz von DNA- und Proteinextraktionen Proteine als Träger des Erbmaterials ausgeschlossen und der Nachweis erbracht werden, dass das Erbmaterial aus DNA besteht (Abb. 78).

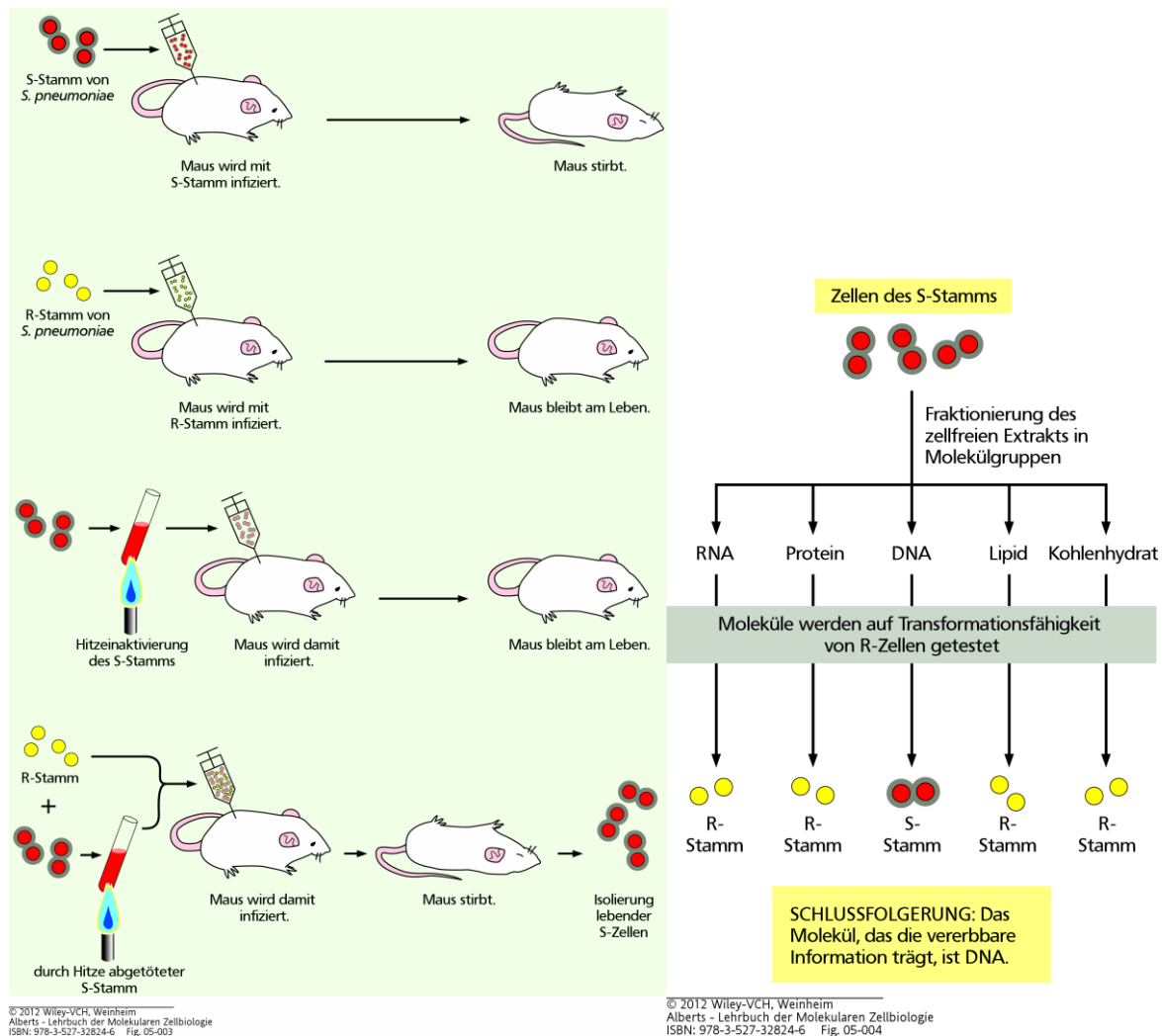


Abbildung 78 Experiment von Griffith, 1928

1952 entdeckten die Forscher Alfred Day Hershey und Martha Chase, dass die genetische Information ausschließlich in DNA enthalten ist.

Die beiden Forscher arbeiteten mit dem Bakteriophagen T2, der das Darmbakterium *Escherichia coli* befällt. Die Bakteriophagen bestehen aus einer Proteinhülle, die man im Experiment mit einem Schwefelisotop (^{35}S) markieren kann. Schwefelatome kommen in Nukleinsäuren nicht vor. DNA kann dagegen mit einem Phosphorisotop (^{32}P) markiert werden. Phosphoratome fehlen wiederum in Proteinen der Bakteriophage. In zwei unabhängigen Experimenten infizierten Hershey und Chase Bakterien mit den T2-Phagen, die entweder radioaktiv markiertes Protein oder radioaktiv markierte DNA enthielten. Es stellte sich heraus, dass nur die radioaktiv markierte DNA in den neu gebildeten Phagen auftauchte, nicht aber die radioaktiv markierten Proteine (Abb. 79).

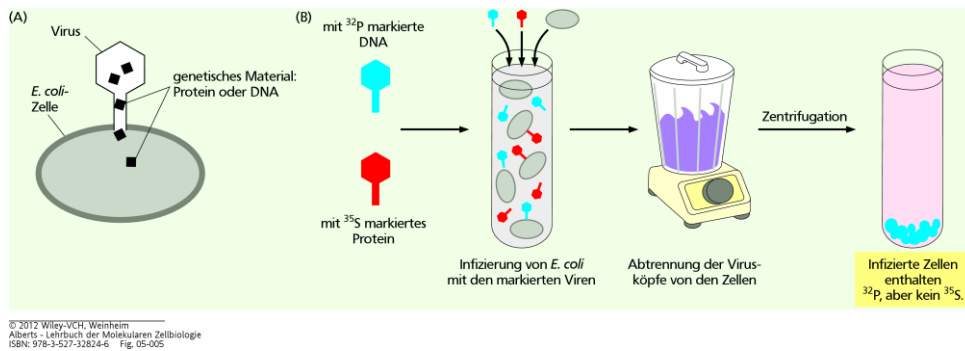


Abbildung 79 Experiment von Hershey & Chase, 1944

Die Experimente von Griffith, Avery, Hershey und Chase gelten heute als die klassischen Arbeiten, deren Ergebnisse die DNA als Träger der genetischen Information identifizierten. Doch woraus besteht die DNA überhaupt?

Nukleinsäuren

Die DNA gehört zur Stoffgruppe der Nukleinsäuren. Es handelt sich hierbei um Polymere, deren Untereinheiten (Monomere) die Nukleotide sind. Nukleotide bestehen aus drei prinzipiellen Bausteinen: Phosphaten, Zucker, Pyrimidin- oder Purin-Basen. Es gibt vier verschiedene Basen in der DNA: Cytosin und Thymin, die zu den Pyrimidin-Basen gehören sowie Guanin und Adenin, die zu den Purin-Basen gehören. Bei dem Zuckermolekül handelt es sich um die 2-Desoxyribose (Abb. 80). Das Adenin haben wir in einer anderen Verbindung, dem ATP, kennengelernt.

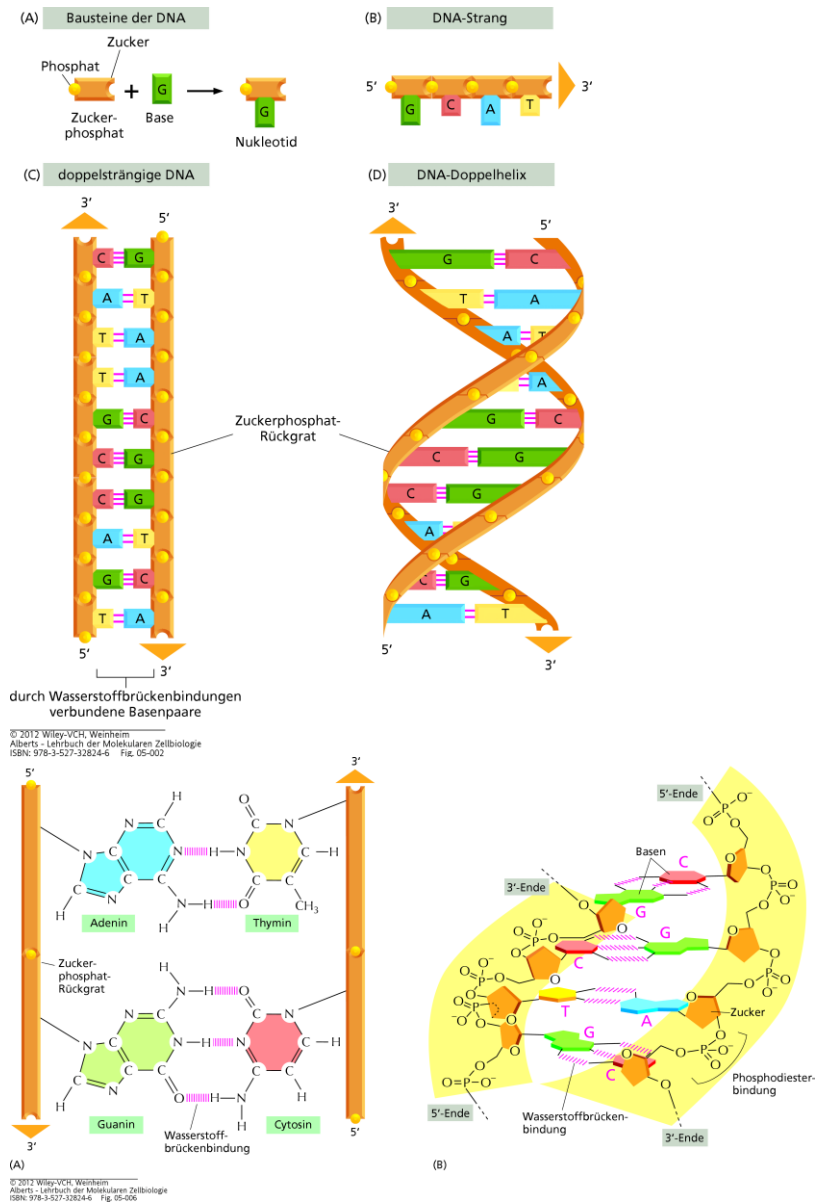


Abbildung 80 Grundaufbau der DNA

Das erste C-Atom des Zuckers Desoxyribose geht eine Verbindung mit einem Stickstoffatom der Base. Für eine Verbindung mit dem Phosphatrest steht das 5 bzw. 3. C-Atom der Desoxyribose bereit.

Die Nukleotide verbinden sich miteinander, indem ein Desoxyribose-Molekül über eine Phosphat-Verbindung mit dem nächsten Desoxyribose-Molekül verknüpft wird. Der Phosphatrest bildet somit eine Bindung mit dem 5. C-Atom des einen Zuckermoleküls und dem 3. C-Atom des nächsten Zuckermoleküls usw. Bei diesen Verbindungen wird Wasser abgespalten, wie z. B. bei der Peptidbindung. Das Rückgrat dieser Nukleotid-Ketten besteht aus alternierenden Zuckermolekülen und Phosphatgruppen. Es resultieren zwei unterschiedliche Enden, welche nach der Lage des Kohlenstoffatoms der Desoxyribose entsprechend als 3' – oder 5'-Ende bezeichnet werden.

So entsteht die Nukleotid-Kette. Dieses Phosphat-Zucker-Gerüst bildet den invariablen Teil der DNA. Variabel sind die Nukleotid-Basen, von denen es wie erwähnt vier verschiedene gibt. Solch eine Nukleotid-Kette bezeichnet man auch als Einzelstrang

Die DNA liegt normalerweise jedoch nicht als Einzelstrang, sondern Doppelstrang vor. Die Basen der Nukleotide sind in der Lage, mit anderen Basen in Wechselwirkung zu treten. Diese Basenpaarungen werden durch Wasserstoffbrücken-Bindungen vermittelt, welche sich zwischen Sauerstoff- bzw. Stickstoffatomen der gegenüberliegenden Basen ausbilden. Einer Purinbase steht in der Regel eine Pyrimidinbase gegenüber. Dabei liegen sich stets ein Cytosin und ein Guanin bzw. ein Thymin und ein Adenin gegenüber. Cytosin und Guanin werden durch drei Wasserstoffbrücken, Thymin und Adenin durch zwei miteinander verbunden. Cytosin und Guanin, sowie Thymin und Adenin sind zueinander komplementär. Die Sequenz der Nukleotidbasen eines Stranges der DNA (in 5'- nach 3'-Richtung) entspricht komplementär der Nukleotidbasensequenz (oder einfach Basensequenz) des anderen Stranges in 3'- nach 5'-Richtung. Die Spezifität der Basenpaarung ist das wichtigste strukturelle Merkmal der DNA. Die Menge von A in einem DNA-Doppelstrang entspricht der von T; ebenso korreliert die Menge von C mit der von G. Diese Beziehung ist auch als Chargaff-Regel bekannt. Der AT- bzw. GC-Gehalt variiert zwischen unterschiedlichen Organismen und wird auch als Klassifizierungsmerkmal herangezogen

1953 hatten James D. Watson und Francis H. Crick die Struktur des DNA-Moleküls herausgefunden, indem sie Daten aus Röntgenbeugungsexperimenten von Maurice Wilkins und Rosalind Franklin verwendeten. In einer am 25. April 1953 in Nature erschienenen Arbeit von einer dreiviertel Seite (A structure for deoxyribonucleic acid, Nature 171: 737–738, 1953) beschreiben Watson und Crick die Struktur der DNA als rechtsgängige Doppelhelix. Die abwechselnden Zuckermoleküle und Phosphatgruppen der Einzelstränge bilden ein außenliegendes Rückgrat, während die Basen in das Helixinnere hinein orientiert sind (Abb. 80). Da sowohl GC- als auch AT-Paarungen jeweils eine Purin- und eine Pyrimidinbase enthalten, haben sie eine sehr ähnliche Raumauffüllung.

Im isolierten Zustand tritt uns die DNA als relativ uniformer Faden mit einem Durchmesser von etwa 2 nm entgegen. Erhöht man die Temperatur einer DNA-haltigen Lösung, werden die Wasserstoffbrücken-Bindungen zerstört und die beiden Stränge trennen sich voneinander zu Einzelstrang-DNA. Mit zunehmender Temperatur entstehen immer mehr einzelsträngige Bereiche. Die Herstellung einzelsträngiger DNA wird auch als Denaturierung bezeichnet. Dies kann durch erhöhte Temperatur oder der Zugabe bestimmter Chemikalien wie Harnstoff oder Formamid erreicht werden.

Bei der Renaturierung, auch als Reassoziatio**n** oder Annealing bezeichnet, lagern sich die komplementären DNA-Stränge wieder aneinander. Dieser Prozess kann durch die Wahl der Pufferbedingungen (pH-Wert bzw. Salzkonzentration), vor allem aber der Temperatur beeinflusst werden.

RNA

Die RNA ist ebenso wie die DNA eine Nukleinsäure. RNA besteht jedoch aus dem Zucker Ribose. Der Unterschied zur Desoxyribose besteht, dass die Ribose am 2. C-Atom eine OH-Gruppe gebunden hat, die 2-Desoxyribose – wie der Name vermuten lässt – hat nur eine Wasserstoffbindung mit dem zweiten C-Atom. Die RNA besitzt an Stelle des Thymins die Pyrimidinbase Uracil (U), die sich vom Thymin nur durch das Fehlen einer Methylgruppe in Position 5 unterscheidet (Abb. 81). Die RNA übernimmt im Zellstoffwechsel, besonders bei der Synthese von Proteinen, eine wichtige Rolle ein.

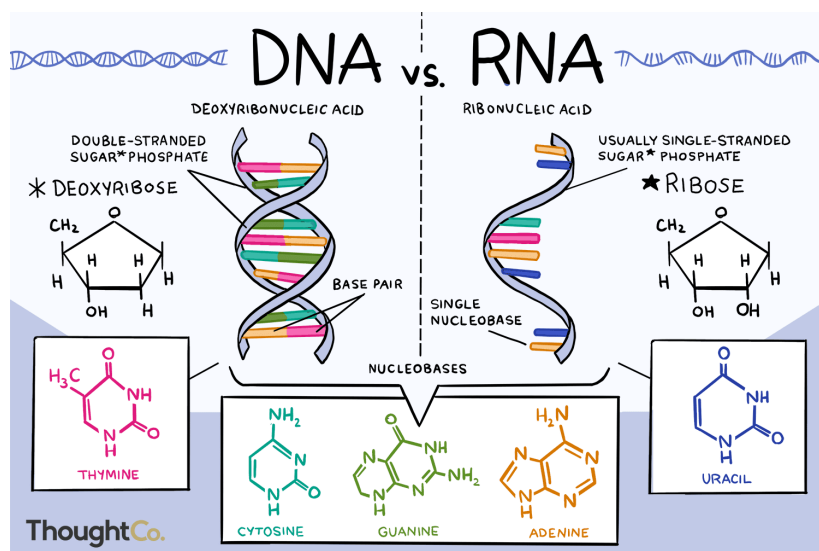


Abbildung 81 RNA

Organisation der Genome

Bei den Bakterien schwimmt die DNA frei im Zellplasma, auch als Bakterien-Chromosom bezeichnet. Dabei ist das Bakterien-Chromosom viel länger als die Bakterienzelle. Bei Coli-Bakterien ist das doppelsträngige DNA-Molekül etwa 1,6 mm lang und damit tausendmal länger als die Bakterienzelle selbst. Daher müssen die Genome kompaktiert werden. dies geschieht mittels Bindeproteinen, die in Wechselwirkung mit dem Bakteriengenom treten und dieses so schleifenförmig organisiert. Dies kann geschehen, weil die DNA negativ geladen ist, die Bindeproteine aber viele positiv geladene Aminosäuren haben (Abb. 82).

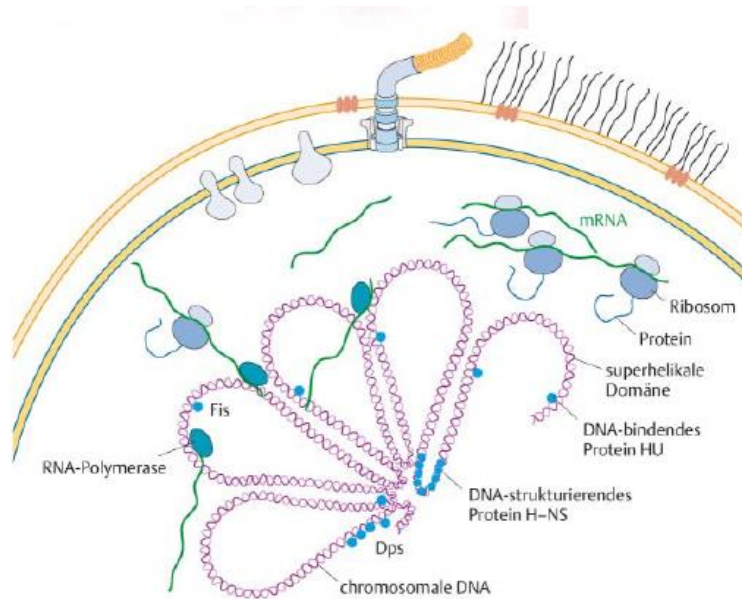


Abbildung 82 Bakteriengenom

Der Großteil des Genoms eukaryotischer Zellen befindet sich im Zellkern. Auch bei Eukaryoten ist der DNA-Faden viel länger als die Zelle und auch hier helfen Bindeproteine das Genom kompakt zu halten. Die DNA und die mit ihr assoziierten Proteine werden zusammen als Chromatin bezeichnet, die kompakteste Form des Genoms als Chromosom. Die wichtigsten Bindeproteine eukaryotischer Zellen bezeichnet man als Histone. Weitere Proteine sind die Nicht-Histon-Proteine, welche für die Organisation von Funktionsdomänen im Chromosom verantwortlich sind.

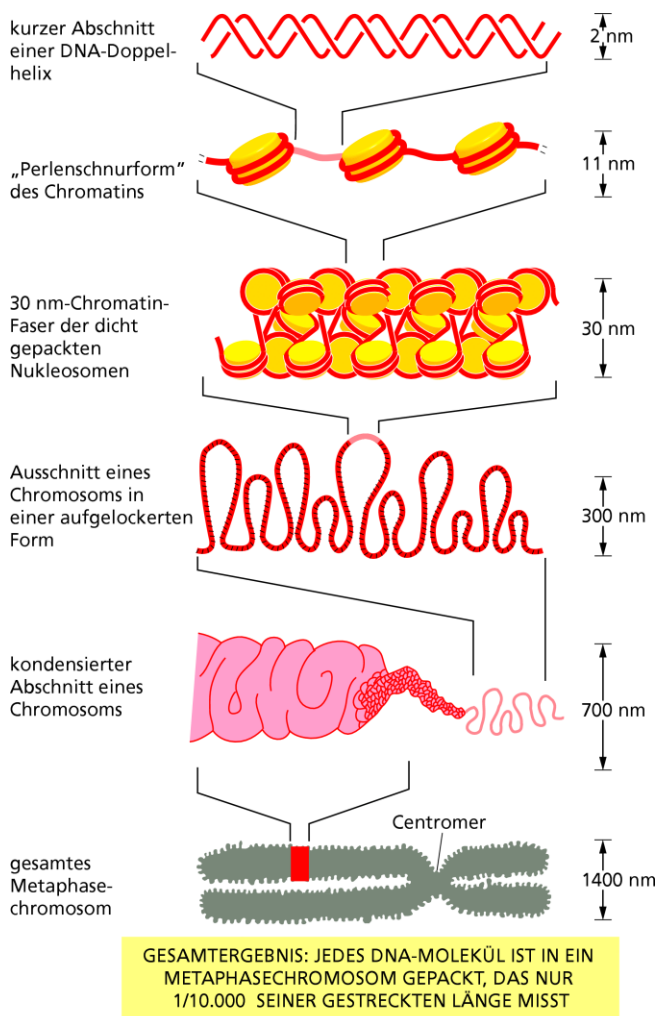
Die Maximale Verkürzung von 1: 12 000 wird über vier Organisationsstufen erreicht (Abb. 83). Die unterste Verpackungsstufe besteht in der Ausbildung der Nukleosomen und führt zu einem 10-nm-Faden. Auf dieser Stufe wird ein Verkürzungsfaktor des DNA-Doppelstrangs von sechs bis sieben erzielt. Ein Nukleosom besteht aus einem bestimmten Abschnitt Kern-DNA, der um ein Aggregat aus acht Histon-Proteinen gewickelt ist. Diese Histone bezeichnet man auch als Core-Histone und das Aggregat als Histon-Oktamer. Es gibt vier verschiedene Core-Histon-Moleküle, die als H2A, H2B, H3 und H4 bezeichnet werden. Alle vier Core-Histone sind im Nukleosom zweimal vertreten und bilden das Histon-Oktamer. Nukleosomen enthalten weiterhin einen außerhalb der Core-Histone liegenden DNA-Abschnitt und ein weiteres Histonmolekül, das Linker-Histon H1.

Die Histone haben zwei auffällige Eigenschaften. Sie besitzen einen relativ hohen Anteil (zusammen über 20 %) an den basischen Aminosäuren Arginin und Lysin, durch deren positive Ladungen die negativen Ladungen der DNA, die den Core-Histonen eng anliegt, neutralisiert werden. Weiterhin sind die vier Core-Histone in ihrer Molekülgröße und Aminosäurezusammensetzung in der Evolution außerordentlich konserviert. Histon H4 unterscheidet sich z. B. beim Rind und bei der Erbse nur durch eine einzige Aminosäure.

Die Sekundärstruktur bildet der sogenannte Solenoid, ein 30 nm-Faden. Bei steigender Ionenstärke in der Lösung bilden die Nukleosomen eine Überstruktur aus. Es ist eine Superhelix mit nur kleinen Unregelmäßigkeiten, bei der das Linker-Histon H1 im Innenraum zu liegen kommt.

Die Tertiärstruktur bilden sogenannte Schleifendomänen aus. Sie werden vom 30-nm-Faden ausgebildet. Sie lagern sich zu mehreren in einer Ebene an (Rosette) und sind in der Mitte über spezifische DNA-Bereiche an ein Proteingerüst gebunden. Sie bilden somit Rosetten. Die Anordnung zahlreicher Rosetten hintereinander führt zu einem 300-nm-Faden. Dieser 300-nm-Faden ist helikal aufgewunden und bildet die Quartärstruktur: die Chromatide. In jeder Chromatidenwindung sind etwa 30 Rosetten enthalten.

Im Zuge der Zellteilung (Mitose), mit der wir uns gesondert befassen werden, werden die Chromatiden verdoppelt. Zwei gleiche Chromatiden bezeichnet man als Schwesterchromatiden und diese bilden zusammen das Chromosom, die stärkste Verkürzung des DNA-Moleküls.



© 2012 Wiley-VCH, Weinheim
 Alberts - Lehrbuch der Molekularen Zellbiologie
 ISBN: 978-3-527-32824-6 Fig. 05-025

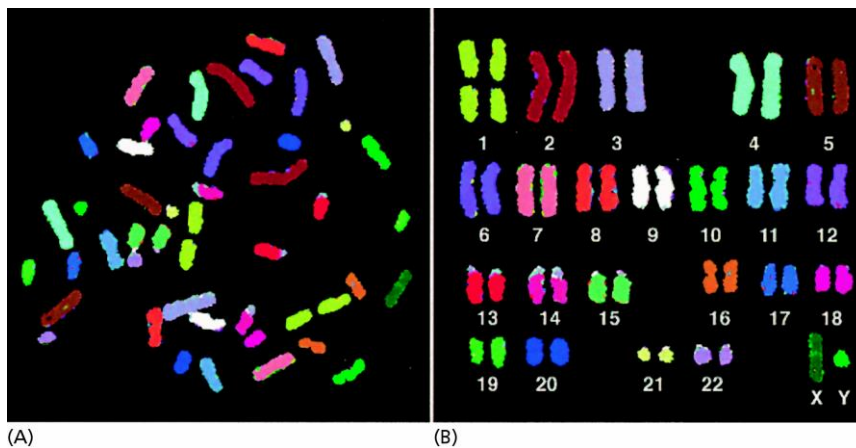
Abbildung 83 Organisationsstufen des eukaryotischen Chromosoms

Chromosomen besitzen eine Reihe von Bereichen, denen spezielle Aufgaben zufallen. Man nennt sie deshalb auch Funktionselemente der Chromosomen.

Zu den wichtigsten gehören die Centromere, die für die geordnete Verteilung der Chromosomen in der Zellteilung verantwortlich sind.

Die Enden der Chromosomen bilden die Telomere, die bei der Verdoppelung der DNA eine wichtige Rolle spielen.

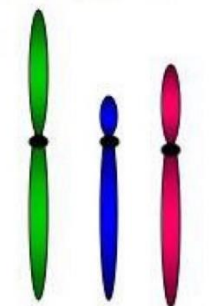
Die Gesamtheit der Chromosomen bezeichnet man als Karyotyp. Jede Tierart verfügt über eine bestimmte Anzahl an Chromosomen. Wir Menschen haben 46 Chromosomen. 44 sogenannte autosomale Chromosomen und 2 Geschlechtschromosomen (Frauen: XX, Männer: XY). Tatsächlich handelt es sich bei den 46 Chromosomen um 23 Chromosomen-Paare. Das heißt: Wir haben 23 unterschiedliche Chromosomen, die jeweils doppelt vorhanden sind. Doppelte Chromosomen sind zueinander homolog (sie sind Variationen desselben Chromosomentyps). Einen doppelten Chromosomensatz bezeichnet man als diploid oder kurz: $2n$. Der Karyotyp kann in einem Karyogramm dargestellt werden. Dieser diploide Chromosomensatz entsteht durch die Verschmelzung der Ei- und Spermazelle, die jeweils nur einen halben Chromosomensatz haben, man spricht von haploid (Abb. 8).



© 2012 Wiley-VCH, Weinheim
Alberts - Lehrbuch der Molekularen Zellbiologie
ISBN: 978-3-527-32824-6 Fig. 05-010

Haploid (n)

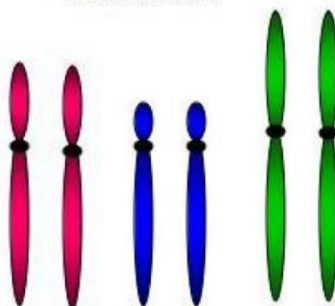
- One copy of genetic material subdivided into chromosomes
- Eg. *Gametes* (sperm and eggs)



Three *nonhomologous* chromosomes

Diploid ($2n$)

- Two copies of genetic material subdivided into chromosomes
- *Somatic cells*



Three pairs of *homologous* chromosomes

Abbildung 84 Karyogramm des Menschen, haploid und diploid

Das menschliche Genom

Die Genomgröße kann bei Lebewesen unterschiedlich groß sein (Abb. 85). Betrachten wir uns aber die DNA des Menschen etwas genauer.

Genomgrößen

Virus-Genom	$5 * 2.000 =$	10.000
Bakterien-Genom	$150 * 2.000 =$	3.000.000
Kleinstes Pflanzengenom (Arabidopsis Thaliana)	$60.000 * 2.000 =$	120.000.000
menschliches Genom	$1.500.000 * 2.000 =$	3.000.000.000
Gersten-Genom	$2.500.000 * 2.000 =$	5.000.000.000
größtes Pflanzengenom	$60.000.000 * 2.000 =$	120.000.000.000



Vorlesung Einführung in die Bioinformatik - Grundlagen

U. Scholz & M. Lange Folie #1-9

Abbildung 85 Genomgrößen

Der DNA-Strang eines Menschen besitzt etwa 3 Milliarden Basenpaare, hat also vereinfacht ausgedrückt 3 Mrd. Buchstaben der entsprechenden Basen A, T, C und G. Nur als Vergleich: Die Bibel, welches nun sicherlich kein dünnes Buch ist, hat etwa 4,4 Millionen Zeichen – Leerzeichen & Sonderzeichen inklusive – das entspricht 0,14% der Basenpaare der DNA.

Aber diese 3 Mrd. Basenpaare können noch weiter unterteilt werden (Abb. 86):

Human genome

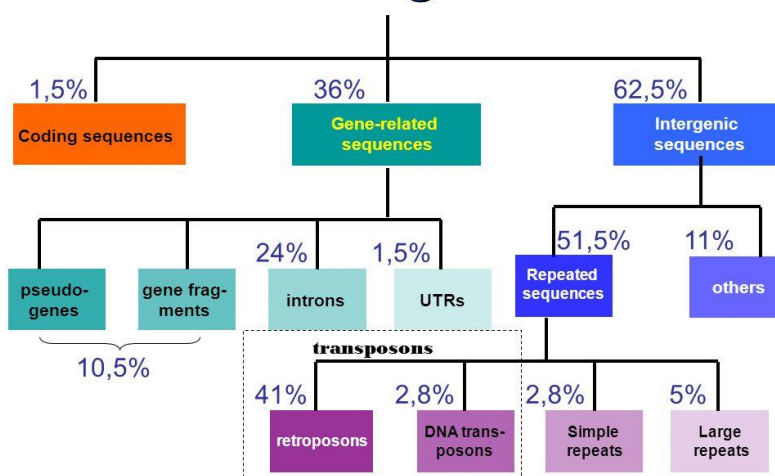


Abbildung 86 menschliches Genom

Zum einen haben wir die Gene. Jeder kennt diesen Begriff bzw. hat von ihm etwas gehört. Generell wird unter Genen das verstanden, was wir biologisch an die nächste Generation vererben. In der Molekularbiologie ist die Definition des Gens eingeschränkter.

Gene definieren sich dadurch, dass sie die Information für den Aufbau von Proteinen haben inklusive ihrer Regulationssequenzen – also jene Abschnitte die mitkontrollieren, wann ein Gen ein- oder ausgeschaltet wird. Solch ein Gen liegt aber oft nicht als „ununterbrochene“ Kette vor, sondern ist „zerstückelt“ in „Introns“ und „Exons“. Exons werden die codierten Elemente genannt, die also die Information für den Aufbau der Proteine haben, die Introns sind die nicht codierenden Bereiche. Über ihre Funktion werden wir uns näher befassen, wenn wir die Proteinsynthese und Genregulation näher behandeln. Das menschliche Genom hat etwa 20.000 Gene. Das klingt viel – beträgt aber nur 1-1,5% der Gesamt-DNA. Der Rest besteht aus sog. nicht-codierter DNA.

Hierzu gehören auch die Pseudogene. Das sind Gene, die ehemals eine Funktion hatten, aber durch Mutationen ausgeschaltet wurden und somit nicht in Proteine übersetzt werden.

Ein Großteil des menschlichen Genoms – und des Genoms aller Eukaryoten – besteht aus sogenannter repetitiver DNA. Repetitive DNA besteht aus sich vielfach wiederholenden DNA-Sequenzen. Die jeweils wiederholten Sequenzen können entweder über das gesamte Genom verstreut oder tandemartig angeordnet an bestimmten Stellen des Genoms liegen (Abb. 87).

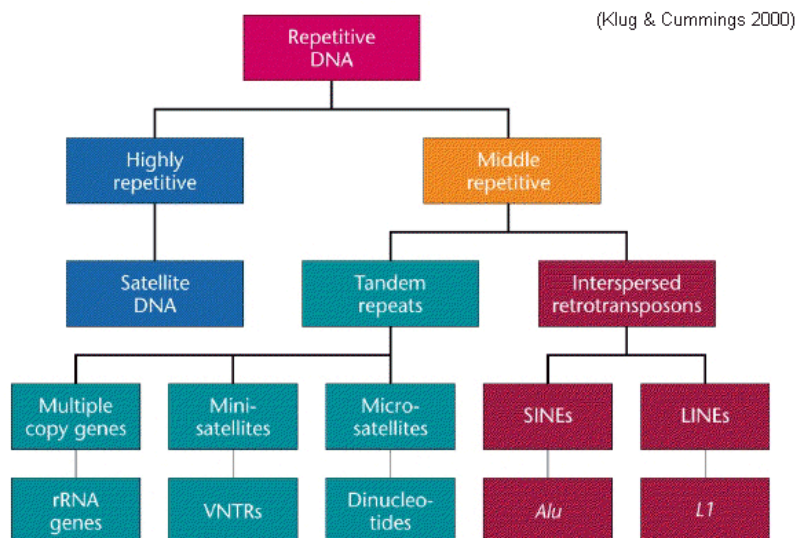


Abbildung 87 Repetitive DNA

Je nach ihrer Größe gibt es klassische Satelliten-DNA, die bis zu 5 Mio. Basenpaare lang sein können und bis zu einer Mio. Wiederholungen von Sequenzen mit 5 bis 300 Nukleotiden haben können. Zu ihnen gehören die Alpha-Sequenzwiederholungen an den Centromeren der Chromosomen, die beim Menschen aus tandemartig hintereinander liegenden Wiederholungen einer Sequenz von 171 bp bestehen und Bindungsstellen für Centromer-spezifische Proteine sind. Mit einer Gesamtlänge von 100 – 20.000 Basenpaaren sind die Gruppen der Minisatelliten-DNA deutlich kleiner.

Zu den Minisatelliten, die aus Einheiten von maximal 15 Basen bestehen, gehört die DNA der Telomere an den Chromosomenenden. Sie enthält bis zu 1000 tandemartige Wiederholungen einer kurzen DNA-Sequenz, die beim Menschen aus 6 Basen besteht. Sie sind Bindungsstellen für Telomer-spezifische Proteine.

Mikrosatelliten sind sehr kurze Einheiten von bis zu 6 Basenpaaren, die nur bis zu einigen hundert Mal tandemartig wiederholt werden. Mikrosatelliten stellen selbst repetitive Einheiten dar, die über alle Chromosomen eines Genoms verteilt bis zu 100 000-mal vorkommen können. Wegen der Zufälligkeit ihrer Entstehung sind Mikrosatelliten bei verschiedenen Individuen unterschiedlich lang (Längenpolymorphismus), was man sich bei der Erstellung eines „genetischen Fingerabdrucks“ zunutze macht. Durch Vergleich der Längen von Genomregionen, die Mini- und Mikrosatelliten-Loci enthalten, lassen sich Individuen eindeutig identifizieren und verwandtschaftliche Beziehungen aufklären.

Mittelrepetitive Sequenzen sind unterschiedlich lang und haben verschiedene Längeneinheiten von Wiederholungen und sind meist über das gesamte Genom verteilt. Hier unterscheidet man zwischen den SINEs (short interspersed elements) und LINEs (long interspersed elements), sie sind über das gesamte Genom verteilt.

SINEs haben kurze Sequenzen von 200 bis 400 Basenpaaren. Zu ihnen gehören die ca. 1,1 Millionen Alu-Elemente, die den größten Teil der mittelrepetitiven Sequenzen des menschlichen Genoms ausmachen.

LINEs haben eine Länge von 1.400 – 6.000 Basenpaaren und kommen mit einer Kopienzahl von 60.000 bis 100.000 im Genom vor.

Ebenfalls den LINEs zugerechnet werden die Transposons, die springenden Gene. Hinzu kommen die menschlichen endogenen Retroviren (HERV, human endogenous retroviruses).

Bei Eukaryoten befindet sich die DNA nicht nur im Zellkern, sondern auch in einigen Zellorganellen, namentlich den Mitochondrien und den Chloroplasten. Eine Zelle kann mehrere hundert bis tausend Mitochondrien haben und jedes Mitochondrium hat einen kleinen DNA-Ring, der bei Menschen 16.569 Basenpaare lang ist.

Kapitel 9: DNA-Replikation

Wir werden alle älter, das ist nicht zu verhindern. Dass wir älter werden, hat z. T. auch mit unserer DNA zu tun, bzw. ihrer Vermehrung, mit der wir uns hier befassen werden.

Wenn sich eine Zelle teilt, so muss ihr Genom auch an die Tochterzellen weitergegeben werden. Damit dies geschehen kann, muss die DNA verdoppelt werden. Als Replikation bezeichnet man einen zellulären Vorgang, bei dem eine DNA-Doppelhelix mithilfe spezifischer Enzyme identisch verdoppelt wird. Nur unter dieser Voraussetzung entstehen bei einer nachfolgenden Zellteilung Tochterzellen mit gleicher genetischer Ausstattung.

Obwohl das Grundschemata der Replikation einleuchtet, sind die Details jedoch sehr komplex. Folgende Besonderheiten sind dabei zu beachten:

1. Die komplementäre Basenpaarung soll fehlerfrei erfolgen: Adenin paart stets mit Thymin und Guanin stets mit Cytosin.
2. Alle bekannten Enzyme, die DNA entlang einer Matrize synthetisieren, arbeiten nur in einer bestimmten Richtung. Außerdem können sie einen vorhandenen Strang zwar verlängern, aber nicht neu beginnen.
3. Die beiden DNA-Stränge in der DNA-Doppelhelix verlaufen antiparallel, d. h. sie haben eine entgegengesetzte Leserichtung. Die Synthese der beiden Tochterstränge wird sich also unterscheiden.
4. Das lange DNA-Molekül liegt im Zellkern verpackt vor.
5. Die Replikation muss zeitlich mit anderen Prozessen an der DNA und in der Zelle abgestimmt sein.

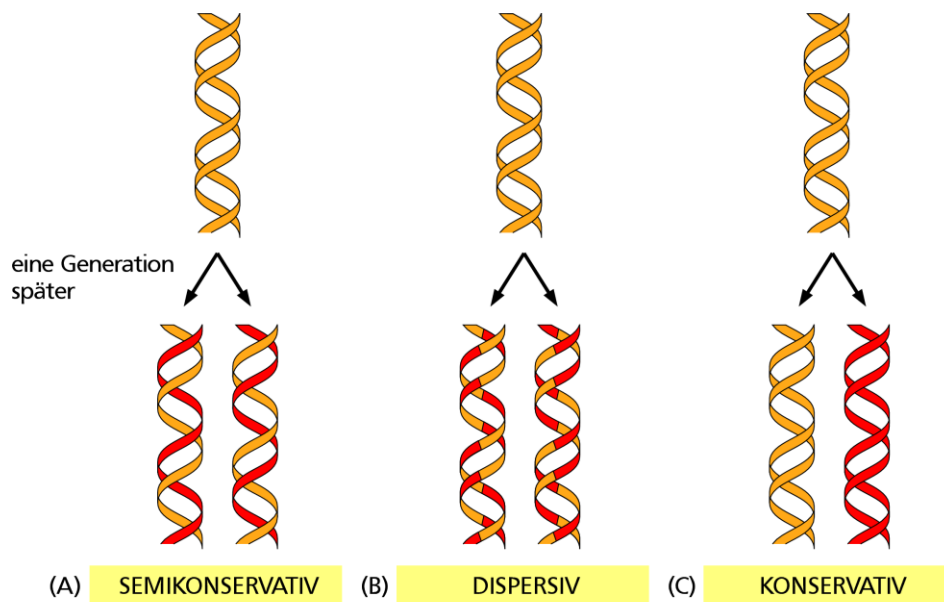
Meselson-Stahl-Experiment

Es gibt theoretisch gesehen drei Möglichkeiten wie die DNA verdoppelt werden könnte (Abb. 88):

Zunächst ist ein Replikationstyp möglich, der als dispersiv bezeichnet wird. Die Tochtermoleküle bestehen aus einer abwechselnden Mischung von alten und neu synthetisierten Strängen.

Von einer konservativen Verdopplung spricht man, wenn eine DNA-Doppelhelix komplett neu synthetisiert wird und die elterliche Doppelhelix unverändert bleibt.

Bei einer semikonservativen Verdopplung trennen sich die Elternstränge und zu jedem der beiden Elternstränge wird der komplementäre Strang neu synthetisiert. So ergeben sich auch hier zwei identische DNA-Doppelhelices; in diesem Fall besteht aber jeder Doppelstrang jeweils aus einem Elternstrang und einem neu synthetisierten Tochterstrang.



© 2012 Wiley-VCH, Weinheim
 Alberts - Lehrbuch der Molekularen Zellbiologie
 ISBN: 978-3-527-32824-6 Fig. 06-006

Abbildung 88 Möglichkeiten der Replikation

Tatsächlich folgt die Replikation der DNA dem semikonservativen Weg. Dies wurde durch das Experiment von Matthew Meselson und Franklin Stahl nachgewiesen (Abb. 89).

Dafür gaben sie die Coli-Bakterien auf ein Nährmedium, dass das Stickstoffisotop ^{15}N enthielt. So kam es zu einem Einbau des Isotops in die DNA der Bakterien. Nach einer Weile überführten sie die Bakterien auf ein Nährmedium mit dem Stickstoffisotop ^{14}N . Die beiden Isotope unterscheiden sich nur im Hinblick auf ihre Masse und das machten sich Meselson und Stahl zu Nutze: Sie extrahierten die DNA der Nachfolgeneration und zentrifugierten das Erbgut. Das Ergebnis (eine Bande) lag genau zwischen den Dichtegradienten von ^{14}N und ^{15}N , was die konservative Replikation ausschloss. Denn sonst hätte das zentrifugierte Erbgut sich auf die Dichtegradienten ^{14}N und ^{15}N aufteilen müssen, weil die Elterngeneration (P1) der Bakterien ^{15}N eingebaut hatte und die erste Tochtergeneration (F1) nur ^{14}N hätte einbauen können. Damit kamen nur noch die semikonservative oder die disperse Replikation in Betracht. Der Versuch wurde um noch eine weitere Tochtergeneration weitergeführt (F2). Doch dieses Mal bekam man zwei verschieden schwere Banden bestehend aus einer leichten ^{14}N Bande und einer schwereren $^{14}\text{N}+^{15}\text{N}$ Bande. Disperse Replikation war damit ebenso ausgeschlossen, denn in diesem Fall hätte sich nur eine Bande bilden dürfen (durch den immer wieder gleichmäßigen Einbau von ^{14}N und ^{15}N).

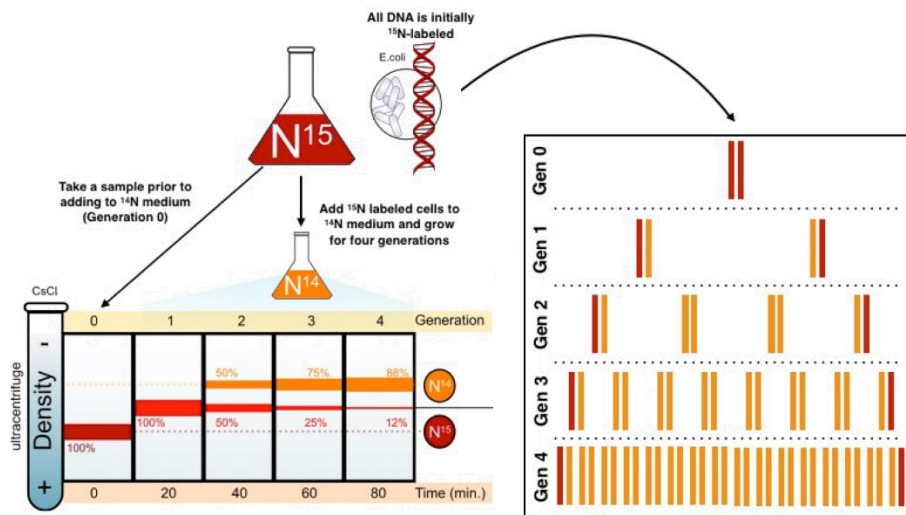


Abbildung 89 Meselson-Stahl-Experiment

Ablauf der Replikation

Wie läuft die Replikation der DNA ab? Hier sind mehrere Enzyme beteiligt (Abb. 90)

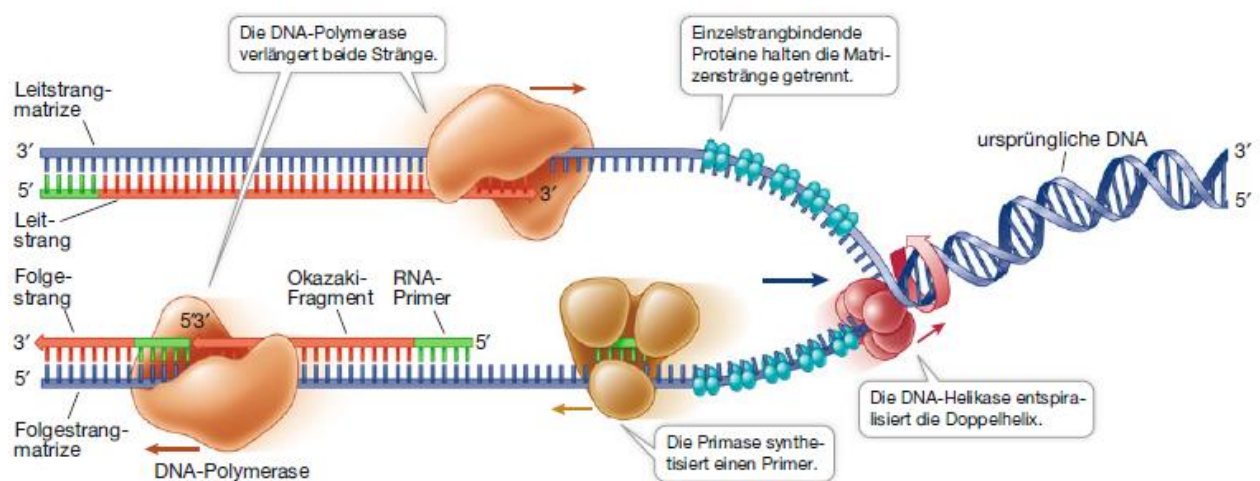


Abbildung 90 Replikation

1. Zuerst muss die Doppelhelix der DNA entwunden werden. Dies macht das Enzym Topoisomerase
2. Bei der Replikation werden die beiden Einzelstränge reißverschlussartig voneinander getrennt, es entsteht somit eine Y-förmige Replikationsgabel. Jeder der frei gewordenen Einzelstränge dient als Matrice für die Synthese des Tochterstrangs. Die Trennung der Doppelstränge wird durch ein Enzym Helikase ermöglicht. Einzelstrangbindende Proteine binden an die getrennten Stränge und verhindern, dass sich dort die Doppelhelix wieder bildet.

3. Die Primase synthetisiert an den 3' Enden der DNA sogenannte Primer, das sind RNA-Moleküle, die für den Beginn der eigentlichen Replikation nötig sind und als Startpunkt dienen.
4. Die Synthese des neuen Tochterstranges erfolgt durch das Enzym DNA-Polymerase III. Und hier stehen wir vor einem kleinen Problem: Die DNA-Polymerase III kann nur in eine Richtung kontinuierlich synthetisieren. Wir haben auf dem Beitrag des Aufbaus der DNA gelernt, dass das Zuckermolekül Desoxyribose mit den 5. und 3. C-Atom eine Verbindung mit dem Phosphat eingeht. Dabei bildet sich eine Phosphatbrücke zwischen dem 5'-C-Atom des einen und dem 3'-C-Atom des anderen Zuckers im benachbarten Nucleotid. Jede Polynukleotidkette hat somit ein 5'-Ende mit einer freien Phosphatgruppe und ein 3'-Ende mit einer freien OH-Gruppe.
5. Die DNA-Doppelhelix ist also aus zwei antiparallelen Strängen aufgebaut, aus einem Vorwärtsstrang in 3'→5'-Richtung und einem Rückwärtsstrang in 5'→3'-Richtung. Am Vorwärtsstrang kann die DNA-Polymerase kontinuierlich den Leitstrang in 5'→3'-Richtung Nukleotid für Nukleotid synthetisieren. Am Rückwärtsstrang wird der Folgestrang synthetisiert. Da er jedoch antiparallel ist, unterliegt er einem besonderen Synthese-Mechanismus.
6. Reiji Okazaki löste dieses Problem im Jahre 1965. Er zeigte experimentell, dass auch der Rückwärtsstrang in 3'→5'-Richtung von der DNA-Polymerase III abgelesen wird – allerdings stückweise. Denn das funktioniert nur, wenn immer wieder neue Primer gesetzt werden. D. h. die Replikationsgabel öffnet sich und an dem Rückwärtsstrang lagern sich versetzt immer neue Primer an. In die Lücken zwischen den Primern kann die DNA-Polymerase III ansetzen und fragmentweise den Folgestrang am Rückwärtsstrang synthetisieren. Diese einzeln synthetisierten Stücke der DNA bezeichnet man als Okazaki-Fragmente. Man spricht auch von einer diskontinuierlichen Bildung des DNA-Stranges.
7. Die DNA-Polymerase I entfernt nun die Primer und ersetzt ihn durch DNA. Zurück bleibt eine kleine Bindungslücke zwischen den aneinandergrenzenden Okazaki-Fragmenten, die von der DNA-Polymerase I nicht geknüpft werden kann. Das Herstellen dieser Bindung wird von der DNA-Ligase katalysiert, wodurch die Fragmente miteinander verbunden werden und der Folgestrang ein Ganzes bildet.
8. Die Replikation der DNA erfordert Energie. Die DNA-Polymerase nutzt als Einzel-Nukleotide zur DNA-Synthese Desoxyribonucleosidtriphosphate. Durch die Abspaltung von zwei Phosphatgruppen wird Energie frei und die nun mit nur einem Phosphat bestückten Einzel-Nukleotide bilden den neuen DNA-Strang.

Telomere und Alterungsprozess

Wir haben im vorherigen Beitrag in Erfahrung gebracht, dass die Enden der Chromosomen der Eukaryoten als Telomere bezeichnet werden. Sie bestehen aus repetitiver DNA. Beim Menschen ist TTAGGG-30 die Telomersequenz, die sich an jedem Ende eines Chromosoms etwa 2500-mal wiederholt. An diese Sequenzwiederholungen binden spezielle Proteine, die verhindern, dass das DNA-Reparatursystem die Enden als Brüche erkennt. Darüber hinaus können die Wiederholungen Schleifen bilden, die eine ähnliche Schutzfunktion besitzen.

Es gibt jedoch noch ein weiteres Problem mit den Enden der Chromosomen. Wie wir erfahren haben, erfolgt die Synthese des Folgestrangs durch das Anfügen von Okazaki-Fragmenten an RNA-Primer. Wenn nun der endständige, letzte RNA-Primer entfernt wird, kann keine DNA mehr synthetisiert werden, um ihn zu ersetzen, da es kein 3'-Ende mehr gibt, das verlängert werden könnte. Es bleibt also an beiden Enden der fertigen DNA am jeweils komplementären Strang ein kurzes, ungepaartes DNA-Stück übrig; meistens wird es enzymatisch entfernt. Auf diese Weise wird das Chromosom bei jeder Zellteilung etwas kürzer. Bei jeder Replikation können Telomere 50 – 200 Basenpaare verlieren.

Nach vielen Teilungen können so die Gene in der Nähe der Chromosomenenden verloren gehen, und die Zelle stirbt. Dieser Effekt erklärt teilweise, warum viele Zelllinien nicht die gesamte Lebenszeit eines Organismus bestehen bleiben: Ihre Telomere gehen verloren. Andererseits verfügen sich ständig teilende Zellen wie die Stammzellen des Knochenmarks und die Urkeimzellen, die die Sperma- und Eizellen produzieren, über einen speziellen Mechanismus, um ihre Telomer-DNA zu erhalten. Das Enzym Telomerase fügt bei diesen Zellen die verloren gegangenen Telomersequenzen an. Die Telomerase enthält eine RNA-Sequenz, die als Matrize für die Sequenzwiederholung der Telomere fungiert.

Bei über 90% aller Tumoren des Menschen wird die Telomerase exprimiert. Wahrscheinlich ist sie ein wichtiger Faktor für die Fähigkeit von Krebszellen, sich ständig zu teilen. Da die meisten normalen Zellen diese Fähigkeit nicht besitzen, ist die Telomerase ein interessantes Ziel für Medikamente, die Tumoren spezifisch angreifen sollen. Auch im Zusammenhang mit dem Alterungsprozess ist die Telomerase im Gespräch. Wenn menschliche Zellkulturen mit einem Gen transformiert werden, das die Telomerase in großen Mengen exprimiert, verkürzen sich deren Telomere nicht. Solche Zellen leben nicht nur 20–30 Generationen und sterben dann, sondern sie werden immortalisiert (unsterblich). Zu klären ist hier noch, inwieweit sich dieser Befund auf das Altern eines ganzen Organismus anwenden lässt.

DNA-Reparatur

Die DNA wird äußerst exakt repliziert und bleibt daher zuverlässig erhalten. Das ist für die Funktionstüchtigkeit aller Zellen essenziell – gleichgültig, ob es sich um einen Prokaryoten oder eine Zelle in einem komplexen, vielzelligen Organismus handelt.

Die Replikation der DNA erfolgt aber nicht mit absoluter Genauigkeit, und DNA unterliegt Veränderungen durch chemische Substanzen und andere Umweltfaktoren.

DNA-Polymerasen machen Fehler bei der Synthese der Polynucleotidstränge – im Allgemeinen wird pro 100.000 replizierten Nucleotiden eine falsche Base eingebaut. Das erscheint zwar angesichts der etwa 3Mrd. Basenpaare beim Menschen vielleicht als wenig bedeutsam, aber die Fehler addieren sich. Wenn die Fehler der DNA-Polymerase nicht repariert würden, gäbe es bei jeder Teilung einer menschlichen Zelle in den neu synthetisierten DNA-Strängen 60.000 falsche Basen. Es ist sogar noch schlimmer: Da die Basen selbst chemisch instabil sind, können sie durch äußere

Faktoren wie Strahlung geschädigt werden. So entstehen Mutationen, an denen die korrekte Paarung verhindert wird.

Vorteilhafterweise werden Fehler bei der DNA-Replikation in den Zellen korrigiert und geschädigte Nucleotide repariert. Zellen verfügen mindestens über drei verschiedene DNA-Reparaturmechanismen (Abb. 91):

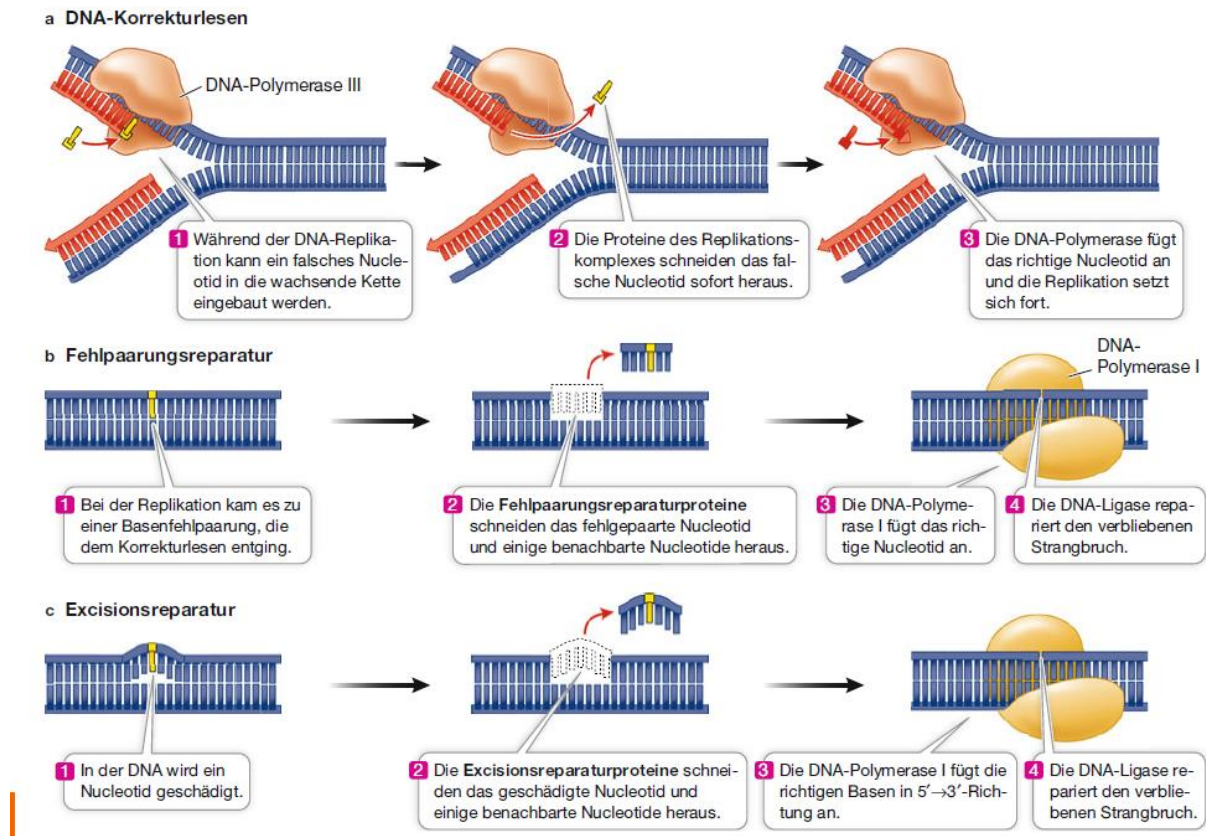


Abbildung 91 Reparatur der DNA

Ein Mechanismus zum Korrekturlesen korrigiert Replikationsfehler, sobald die Polymerase sie erzeugt. Die meisten DNA-Polymerasen führen die Korrekturlesefunktion jedes Mal durch, wenn sie ein neues Nucleotid in den wachsenden DNA-Strang einfügen. Wenn eine DNA-Polymerase eine Fehlpaarung der Basen erkennt, entfernt sie das unpassend eingebaute Nucleotid und startet einen neuen „Versuch“.

Die Fehlerrate dieses Vorgangs beträgt nur eine von 10.000 reparierten Basenpaaren.

Dadurch verringert sich die Fehlerrate bei der Replikation insgesamt auf einen Fehler pro 10Mrd. replizierten Basen.

Ein zweiter Mechanismus zur Fehlpaarungsreparatur prüft die DNA unmittelbar nach der Replikation und korrigiert alle Basenfehlpaarungen (mismatches).

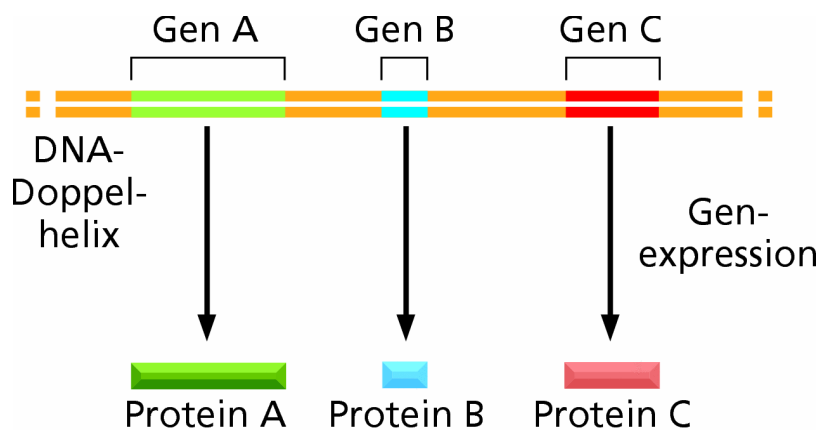
Nach der Replikation der DNA durchsucht eine zweite Gruppe von Proteinen das neu replizierte Molekül auf Basenfehlpaarungen, die dem Korrekturlesen entgangen sind. So kann ein solcher Mechanismus zur Fehlpaarungsreparatur beispielsweise ein A-C-

Paar erkennen, wo eigentlich ein A-T-Paar sein sollte. Das Reparatursystem erkennt, welche der beiden Basen im A-C-Paar die falsche ist und führt die Reparatur durch. Wenn das korrekte Paar A-T ist, ersetzt das Reparatursystem C durch T und stellt so das A-T-Paar wieder her.

DNA-Moleküle können auch während der Interphase einer Zelle beschädigt werden, wenn sich also die DNA nicht verdoppelt. Hochenergetische Strahlung, chemische Substanzen aus der Umgebung und spontan auftretende chemische Reaktionen können die DNA schädigen. Wenn beispielsweise benachbarte Thyminе auf demselben DNA-Strang ultraviolettes Licht absorbieren, bilden sich kovalente Bindungen zwischen den Basen, sodass ein Thymindimer entsteht. Diese Dimere stören die Basenpaarung bei der Replikation, sodass zufällige Basen eingebaut werden. Dies ist die primäre Ursache für Hautkrebs bei Menschen. Für diese Art von Schäden sind Mechanismen der Excisionsreparatur als dritte Möglichkeit, zuständig.

Kapitel 10: Von der DNA zum Protein

In unserem Kapitel über den Aufbau der DNA stellten wir fest, dass diese die Informationen für den Aufbau von Proteinen enthält, die Gene. Proteine haben wir ebenfalls kennengelernt. Proteine können quasi gesehen als Genprodukte verstanden werden. Werden Gene abgelesen und in Proteine übersetzt spricht man auch von Genexpression. Das zugehörige Verb lautet exprimieren (Abb. 92). Die Gesamtheit der Gene wird als Genotyp bezeichnet, die Merkmale für die sie codieren – als Phänotyp.



© 2012 Wiley-VCH, Weinheim
Alberts - Lehrbuch der Molekularen Zellbiologie
ISBN: 978-3-527-32824-6 Fig. 05-009

Abbildung 92 Genexpression

Ein-Gen-Ein-Enzym-(Polypeptid)-Hypothese

Schon in den 1940er Jahren führten George Beadle und Edward Tatum Untersuchungen durch, um die Phänotypen von *Neurospora* chemisch zu bestimmen (Abb. 93). *Neurospora* gehört zu den Pilzen, die die meiste Zeit ihres Lebens haploid sind. Beadle und Tatum stellten die Hypothese auf, dass die Expression eines spezifischen Gens zur Aktivität eines spezifischen Enzyms führt. Nun wollten sie diese Hypothese direkt testen. Sie ließen *Neurospora* auf einem Nährmedium wachsen, das Saccharose, Mineralsalze und Biotin enthielt, das einzige Vitamin, das die Wildtypform von *Neurospora* nicht selbst produzieren kann. Auf diesem Minimalmedium konnten die Enzyme des *Neurospora*-Wildtyps alle Stoffwechselreaktionen katalysieren, die für das Wachstum erforderlich sind.

Beadle and Tatum Experiments				
Bread Mold	Minimal Medium (MM)	MM + Ornithine	MM + Citrulline	MM + Arginine
Wild type	grew	grew	grew	grew
Mutant 1	did not grow	grew	grew	grew
Mutant 2	did not grow	did not grow	grew	grew
Mutant 3	did not grow	did not grow	did not grow	grew

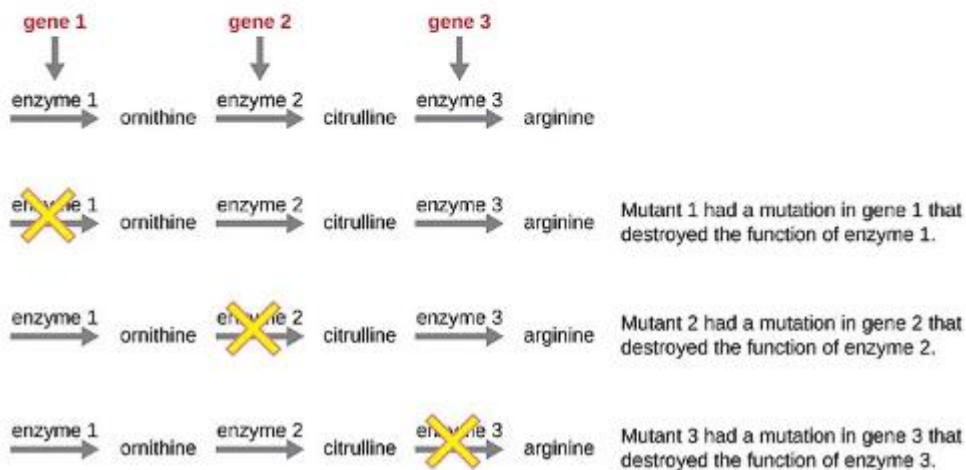


Abbildung 93 Experiment von Beadle & Tatum

Die Wissenschaftler behandelten die Wildtypform von *Neurospora* mit Röntgenstrahlen, die als Mutagen wirken. Ein Mutagen ist ein Faktor, der DNA schädigt und Mutationen hervorruft. Nach der Behandlung mit Röntgenstrahlen konnten einige *Neurospora*-Stämme nicht mehr auf Minimalmedium wachsen. Diese mutierten Stämme wuchsen nur noch, wenn man dem Medium spezifische Nährstoffe zusetzte, etwa bestimmte Vitamine.

Beadle und Tatum stellten die Hypothese auf, dass diese genetischen Stämme in den Genen, die für die Produktion von Enzymen verantwortlich sind Mutationen trugen. Jede Mutation führte in einem bestimmten Gen zur Funktionslosigkeit des Enzyms. Dies ist die sogenannte Ein-Gen-ein-Enzym-Hypothese, bzw. erweitert als Ein-Gen-ein-Protein-Hypothese.

Die Ein-Gen-ein-Protein-Hypothese hat mit zunehmenden Erkenntnissen der molekularen Biologie mehrere Erweiterungen erfahren. Viele Proteine, darunter auch zahlreiche Enzyme, bestehen aus mehr als einer Polypeptidkette oder Untereinheit, das heißt, sie besitzen eine Quartärstruktur. Es ist also korrekter, von der Beziehung Ein-Gen-ein-Polypeptid zu sprechen.

Die Funktion eines Gens besteht darin, die Information für die Produktion eines einzigen spezifischen Polypeptids zu liefern.

Genexpression und zentrales Dogma der Molekularbiologie

Doch die Gene werden nicht direkt in Protein übersetzt, sondern über einen Zwischenweg, die Bildung von RNA. Die RNA haben wir ebenfalls kennengelernt. Sie ist eine Nukleinsäure wie die DNA, hat aber als Zucker die Ribose und statt der Base Thymin die Base Uracil.

Werden die Informationen aus der DNA in Proteine „übersetzt“, spricht man von der Eiweißsynthese oder Proteinbiosynthese. Diese Verläuft in zwei Teilprozessen, Transkription und Translation, die wir hier kennenlernen werden.

Aus dieser Erkenntnis lässt sich das sog. „zentrale Dogma“ der Molekularbiologie schließen (Abb. 94):

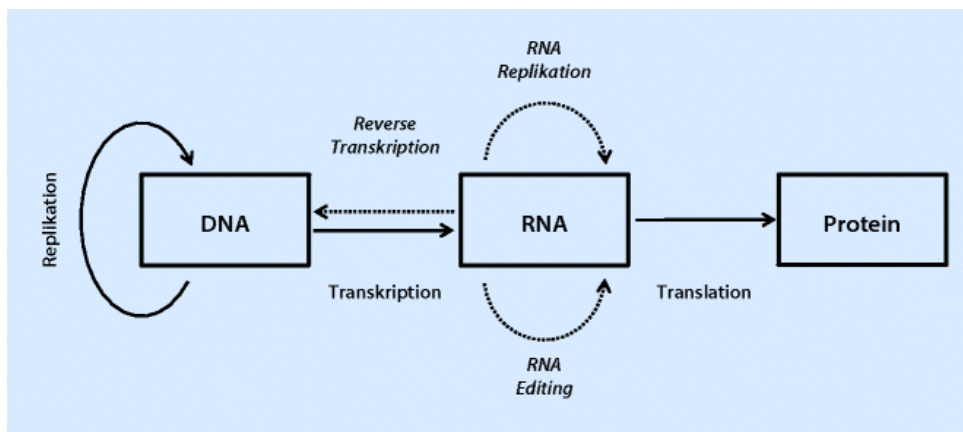


Abbildung 94 Zentrales Dogma der Molekularbiologie

Der Informationsfluss fließt von der DNA zur RNA zum Protein. Der Begriff „zentrales Dogma“ ist unglücklich gewählt, wirkt es doch fast nach Glaubenssätzen der katholischen Kirche. Aber im Prinzip trifft diese Informationsrichtung in der Proteinbiosynthese zu und wir werden sehen, dass sich aus der Struktur des Proteins nicht die DNA-Sequenz genau ableiten lässt. Aber mittlerweile hat dieses „zentrale Dogma“ einige Revidierungen erfahren. So ist bei manchen Viren bekannt, dass die RNA in DNA umwandeln können, wenn sie das Enzym Reverse Transkriptase besitzen. Das trifft z. B. auf Retroviren wie das HIV zu. RNA-Viren können auch ohne DNA Proteine bilden, was wir von Corona-Viren wissen. Und mittlerweile ist auch bekannt, dass einige Proteine die Information zur Bildung von Proteinen haben. Das ist bei Prionen bekannt. Prionen, die eine spezielle Faltung im Vergleich zum ursprünglichen Normalprotein haben, sind in der Lage andere Proteine in Prionen umzuwandeln. Rinderwahn und die Creutzfeldt-Jakob-Krankheit werden von Prionen verursacht. Proteine haben zwar nicht die Information DNA zu bilden, können aber mittels Genregulation Einfluss auf die Genexpression nehmen.

Transkription

Damit ein Protein-kodierendes Gen exprimiert werden kann, muss es zuerst transkribiert werden. Bei der Transkription wird der Code in der DNA des Gens in einen komplementären Code in einem RNA-Molekül umgewandelt. Diese RNA wird als Boten-RNA, Messenger-RNA oder kurz mRNA bezeichnet. Die Transkription erfolgt in drei Schritten: Initiation, Elongation und Termination (Abb. 95).

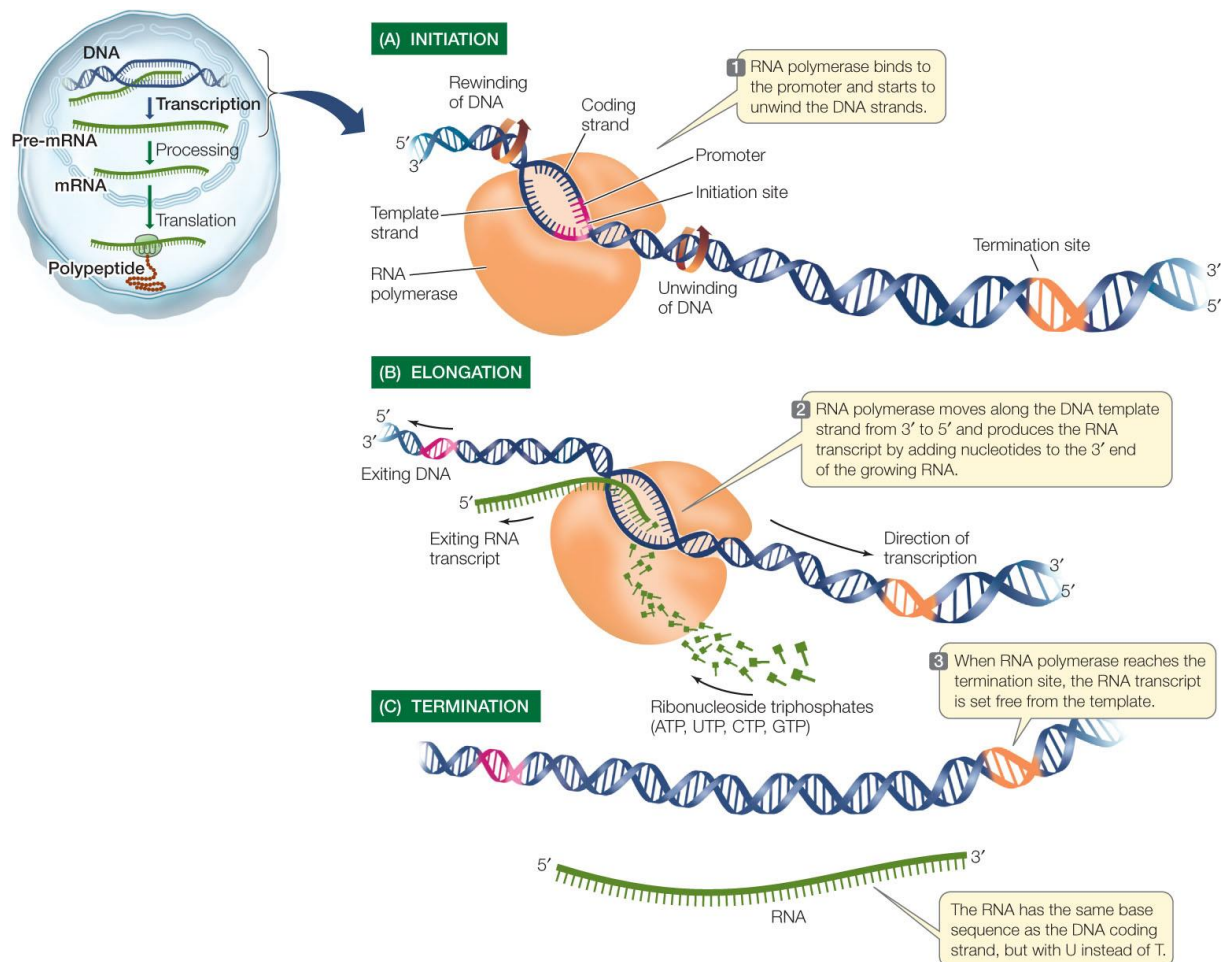


Abbildung 95 Transkription

Initiation

Auch für die Transkription braucht es Enzyme. Hier ist das Enzym RNA-Polymerase entscheidend. Die Initiation der Transkription erfolgt an einem Promotor, einer Region auf der DNA, an der die RNA-Polymerase bindet. Die RNA-Polymerase bindet an den Promotor und beginnt, die DNA abzuwickeln. Nun liegt der DNA-Doppelstrang offen und die Information für die Synthese liegt frei. Aber welcher der beiden DNA-Stränge wird abgelesen, dient also als Matrize für die RNA-Synthese? Wie die DNA-Polymerasen nutzen auch die RNA-Polymerasen den 3'-5'-Strang der DNA als Matrize zur Synthese von mRNA, welche dann in 5'-3'-Richtung gebildet wird. Diesen Strang bezeichnet man auch als Anti-Sense-Strang oder codogener Strang. Der andere DNA-Strang (5'-3') ist der Sense-Strang oder codierender Strang und hat die Basensequenz des entsprechenden Gens.

Mit dem Anti-Sense-Strang als Matrize für die RNA-Synthese entsteht eine RNA, deren Nucleotidsequenz identisch ist mit der des DNA-Sense-Strangs. Die RNA wird also komplementär zu ihrer DNA-Matrize synthetisiert und hat dadurch die gleiche Sequenz wie der Sense-Strang, also um das eigentliche Gen. Damit ist die betreffende DNA-Information direkt in der RNA gespeichert (Abb. 96).

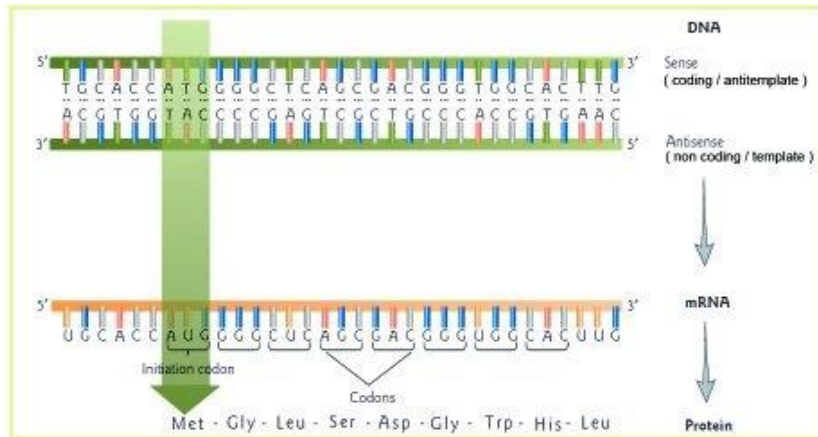


Abbildung 96 Sense – und Anti-Sense-Strang der DNA

Im Gegensatz zu DNA-Polymerasen benötigen RNA-Polymerasen keinen Primer.

Teil eines jeden Promotors ist die Initiationsstelle, an der die Transkription beginnt. Nucleotidgruppen, die „stromaufwärts“ der Initiationsstelle liegen, unterstützen die RNA-Polymerase bei der Bindung.

Weitere Proteine, die an spezifische DNA-Sequenzen und an die RNA-Polymerase binden können, tragen dazu bei, die RNA-Polymerase zum Promotor zu dirigieren. Diese Proteine, die man bei den Prokaryoten als Sigmafaktoren und bei den Eukaryoten als Transkriptionsfaktoren bezeichnet, haben einen Einfluss darauf, welche spezifischen Gene in einer Zelle zu einem bestimmten Zeitpunkt exprimiert werden.

Jedes Gen besitzt zwar einen Promotor, aber nicht alle Promotoren sind identisch. Einige Promotoren sind bei der Initiation der Transkription effizienter als andere. Transkriptionsfaktoren und Promotoren-Effizienz werden wir bei der Genregulation näher kennenlernen.

Elongation:

Sobald die RNA-Polymerase an den Promotor gebunden hat, beginnt sie mit der Elongation. Die DNA wird bei jedem Schritt über einen Abschnitt von etwa zehn Basenpaaren entwunden, und die RNA-Polymerase liest den Anti-Sense-Strang (er verläuft in 3'-5'-Richtung) ab. Die neue RNA wird beginnend mit der ersten Base, die das 5'-Ende bildet, durch Anfügen von Nucleotiden, die zur DNA-Sequenz komplementär sind, zum 3'-Ende hin verlängert. Dieser Prozess kostet Energie. Bei den Einzel-Nukleotiden für die mRNA, die die RNA-Polymerase nutzt, handelt es sich um Ribonucleosidtriphosphate: Durch die Abspaltung von zwei Phosphatgruppen wird Energie frei und die nun mit nur einem Phosphat bestückten Einzel-Nukleotide bilden die mRNA.

Termination:

Genauso wie Initiationsstellen im DNA-Matrizenstrang den Startpunkt der Transkription festlegen, bestimmen auch bestimmte Basensequenzen ihre Termination. Es gibt zwei Mechanismen für die Beendigung der Transkription. Bei manchen Genen faltet sich das neu entstandene Transkript auf sich selbst zurück und bildet interne Wasserstoffbrücken zwischen den Basen. So entsteht eine Schleife, durch die das neu gebildete Transkript von der DNA-Matrize und der RNA-Polymerase abfällt. In anderen Fällen bindet ein Protein an spezifische Sequenzen auf dem Transkript und bewirkt so, dass sich das Transkript von der DNA-Matrize löst.

Genetischer Code

Der genetische Code ist der Schlüssel, mit dessen Hilfe die Nucleotidsequenz einer mRNA, die einem Gen entspricht, in eine Aminosäuresequenz, aus der das von dem Gen codierte Protein besteht, übersetzt (translatiert) wird (Abb. 97). Das heißt, der genetische Code legt fest, welche Aminosäuren nacheinander verwendet werden, um eine bestimmte Polypeptidkette zu bilden.

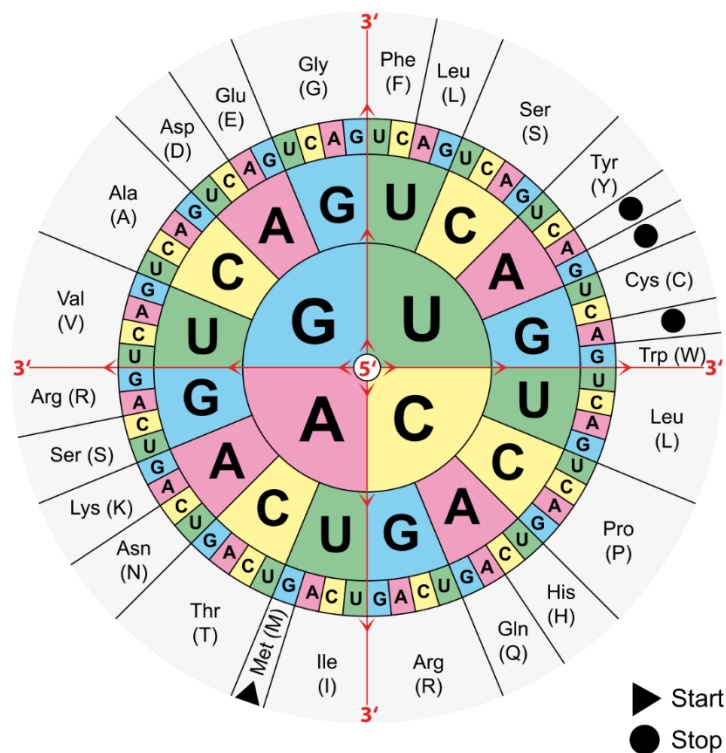


Abbildung 97 genetischer Code

Man kann sich den genetischen Code als eine spezielle Folge nicht-überlappender Buchstabenkombinationen vorstellen. Dabei werden die einzelnen Basen der RNA: Adenin, Cytosin, Uracil und Guanin durch die Buchstaben A, C, U und G abgekürzt. Der genetische Code der RNA ist komplementär zum Anti-Sense-Strang der DNA, der als Matrize diente. Damit ist der Code der RNA identisch zur Basensequenz des Sense-Stranges, der die Information für das Gen (die Aminosäurenabfolge) hat. Drei

einzelne Nukleotide hintereinander bilden einen Code, quasi die Information für eine Aminosäure. So steht z. B. die Kombination der Buchstaben AAA (3x Adenin) für die Aminosäure Lysin. Der genetische Code ist in den 1960er Jahren entschlüsselt worden. Die damalige Fragestellung war, Wie lassen sich mit einem Alphabet, das nur aus vier Buchstaben besteht, mindestens 20 Codewörter schreiben? Oder anders ausgedrückt, wie können die vier RNA-Basen (A, U, G und C) 20 verschiedene Aminosäuren codieren? Da es nur vier Buchstaben gibt, hätte ein Ein-Buchstaben-Code nur vier Aminosäuren codieren können, ein Zwei-Buchstaben-Code nur $4 * 4 = 16$. Aber ein Tripletcode mit $4*4*4 = 64$ Codons würde mehr als genügen, um 20 Aminosäuren zu codieren.

Nun haben wir es aber so, dass wir 64 Codes haben, aber nur 20 Aminosäuren brauchen. Das führt dazu, dass manche Aminosäuren von mehr als nur einem Code bestehen. Der genetische Code ist somit redundant (manche sagen auch „degeneriert“). Die erwähnte Aminosäure Lysin wird z. B. nicht nur durch die Kombination aus AAA codiert, sondern auch die Kombination AAG. Leucin hat sogar 6 Codes.

Daraus ist übrigens ersichtlich, weshalb eine Rückführung der Aminosäuresequenz zur DNA-Sequenz nicht so einfach ist, da mehrere Aminosäuren für mehrere Basensequenzen stehen können.

Von den 64 Codes stehen 61 für Aminosäuren zur Verfügung. Wobei die Aminosäure Methionin durch den Code AUG codiert wird und als Startcodon für die Initiation der Transkription dient. D. h. jede Proteinsynthese startet mit Methionin als Erkennungssequenz für das Ablesen der Gene. Die drei weiteren Codes (UAA, UAG und UGA) sind sog. Stopcodons (Terminationssignale).

Fast alle Spezies auf diesem Planeten verwenden denselben genetischen Code. Dieser Code wurde in der Evolution des Lebens nahezu durchweg vollständig beibehalten und muss daher immens alt sein. Einige Ausnahmen sind jedoch bekannt: So weicht beispielsweise der Code der mitochondrialen DNA und der Chloroplasten-DNA geringfügig von dem Code der Prokaryoten und der Zellkern-DNA der eukaryotischen Zellen ab. UAA und UAG sind dort keine Stoppcodons, sondern codieren Glutamin. Die Bedeutung dieser Unterschiede ist noch unbekannt. Auf jeden Fall ist die Zahl der bekannten Ausnahmen sehr gering.

Marshall Nirenberg und Heinrich Matthaei schafften 1961 den ersten Durchbruch zur Entschlüsselung des genetischen Codes, als sie anstelle der komplexen natürlichen mRNA ein einfaches, synthetisches Polynucleotid als Informationsträger verwenden konnten. Sie begannen mRNA-Moleküle zu synthetisieren, die nur aus einer einzigen Art von Nucleotiden bestanden. So besteht eine Poly(U)-mRNA nur aus Uracilnucleotiden.

Nirenberg und Matthaei wollten nun durch eine Translation im Testgefäß das Polypeptid ermitteln, das diese künstliche mRNA codierte. Ihr Experiment führte zur Identifizierung der ersten Codons. Andere Wissenschaftler fanden bald darauf auch den übrigen Code heraus.

Prozessierung & Splicing

Da alle Lebewesen den genetischen Code gemeinsam haben, könnte man jetzt annehmen, dass der Vorgang der Genexpression bei Eukaryoten und Prokaryoten sehr stark übereinstimmt.

Für den grundlegenden Mechanismus trifft das auch zu, in den Details gibt es jedoch Unterschiede, die wir hier kennenlernen werden.

Die Sequenz einer mRNA, die in einer Zelle an einem Ribosom in ein Polypeptid übersetzt wird, entspricht sowohl bei Prokaryoten als auch bei Eukaryoten einer bestimmten Gensequenz in der DNA dieses Lebewesens. Bei den Eukaryoten befinden sich aber Abschnitte innerhalb eines Gens, welche nicht für die Proteinsynthese verwendet werden.

Gene bei Eukaryoten liegen oft nicht als „ununterbrochene“ Kette vor, sondern sind „zerstückelt“ in „Introns“ und „Exons“. Exons werden die codierten Elemente genannt, die also die Information für den Aufbau der Proteine haben, die Introns sind die nicht codierenden Bereiche (Abb. 98).

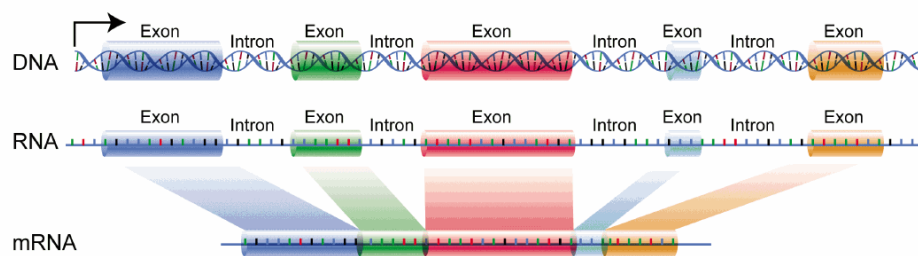


Abbildung 98 Introns & Exons

Die nichtcodierenden Introns werden zwar transkribiert – es entsteht erst mal eine prä-mRNA mit den transkribierten Introns und Exons, aber dann im Zellkern aus der Prä-mRNA entfernt. Nur die Exons verbleiben so schließlich in der bearbeiteten mRNA, die den Zellkern verlässt. Den Schritt, in dem die Introns entfernt werden, bezeichnet man als RNA-Spleißen. Das RNA-Spleißen erfolgt durch das Spleißosom. Dieser RNA-Protein-Komplex entfernt die Introns und verknüpft die Exons (Abb. 99).

Sobald eine Prä-mRNA transkribiert wird, wird sie schnell von mehreren Komplexen gebunden, die als Kleine Kern-Ribonukleoprotein-Partikel (small nuclear ribonucleoprotein particles) oder kurz snRNPs (ausgesprochen „Snurps“) bezeichnet werden. Wie der Name schon sagt, enthalten snRNPs RNAs und Proteine. Die snRNPs sind für das Spleißen von Introns aus Prä-mRNAs verantwortlich.

snRNPs binden an Stellen in einer Prä-mRNA an oder nahe den Intron-Exon-Grenzen. Diese als Konsensussequenzen bezeichneten Stellen enthalten Nukleotidsequenzen, die von den meisten Prä-mRNAs gemeinsam genutzt werden. Die snRNPs enthalten RNA-Moleküle, die durch komplementäre Basenpaarung an diese Konsensussequenzen binden können.

Die snRNPs bilden schließlich einen großen Komplex, der als Spleißosom bezeichnet wird. Während sich das Spleißosom bildet, tritt das Intron aus. Das Spleißosom schneidet die Prä-mRNA an einer Intron-Exon-Grenze, wo es eine reaktive freie Hydroxylgruppe (-OH) am Exon hinterlässt. Das Spleißosom verwendet diese Hydroxylgruppe, um das andere Ende des Introns anzugreifen, und entfernt dabei das Intron und verbindet die Exons miteinander, wobei ein reifes mRNA-Molekül gebildet wird.

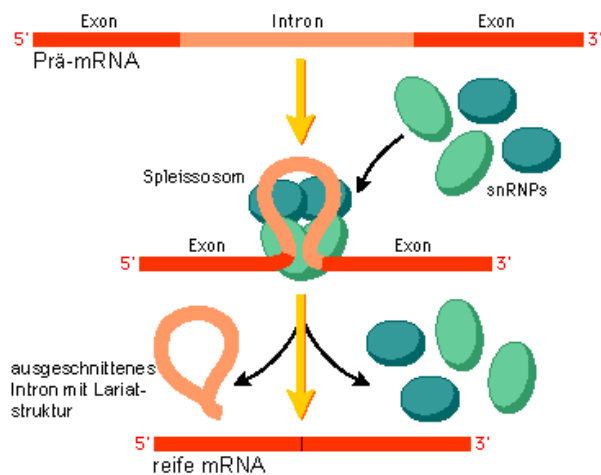


Abbildung 99 Splicing & Spleißosom

Neben dem Spleißen gibt es noch weitere Veränderungen der mRNA. Denn diese muss ja ungehindert den Zellkern verlassen und an die Ribosomen im Zellplasma andocken, ohne abgebaut zu werden (Abb. 100).

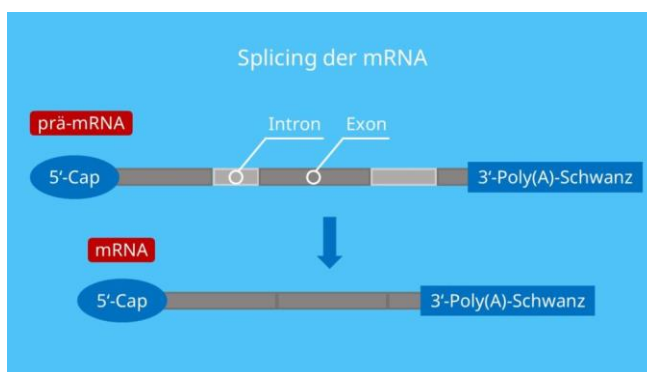


Abbildung 100 Prozessierung der mRNA

Am 5'-Ende der Prä-mRNA wird während der Transkription eine 5'-Cap-Gruppe angehängt. Dies ist ein chemisch modifiziertes Guanosintriphosphat (GTP). Dadurch wird später bei der Translation die Bindung der mRNA an das Ribosom unterstützt und die mRNA ist vor einem Abbau geschützt.

Am 3'-Ende der Prä-mRNA wird ein Poly(A)-Schwanz angehängt. Die Transkription eines Gens endet stromabwärts des Terminationscodons in der DNA. Bei den Eukaryoten liegt normalerweise hinter dem letzten Codon in der Nähe des 3'-Endes der Prä-mRNA eine sogenannte Polyadenylierungs-Sequenz (AAUAAA). Diese Sequenz wirkt als Signal auf ein Enzym, das die Prä-mRNA abschneidet. Unmittelbar nach diesem Schnitt hängt ein anderes Enzym 100–300 Adenin-Nucleotide an das 3'-

Ende der Prä-mRNA. Dieser Poly(A)-Schwanz unterstützt den Export der mRNA aus dem Zellkern und ist für die Stabilität der mRNA von Bedeutung.

Translation

Die Transkription und posttranskriptionelle Vorgänge bringen eine mRNA hervor, die in eine Aminosäuresequenz translatiert werden kann. Die Information der mRNA wird in eine Aminosäuresequenz übersetzt. In der Biologie ist dieser Übersetzer eine besondere Art von RNA-Molekül, das man als Transfer-RNA (tRNA) bezeichnet. Um eine genaue Übersetzung sicherzustellen – das heißt, das Protein wird genau nach den Vorgaben der mRNA synthetisiert – müssen die tRNAs erstens jedes Codon korrekt ablesen und zweitens die zugehörigen Aminosäuren zum Codon bringen. Die Translation findet im Cytoplasma an den Ribosomen statt, sie werden häufig als Proteinfabriken bezeichnet.

Für jede der 20 Aminosäuren gibt es mindestens ein spezifisches tRNA-Molekül. Jede tRNA besitzt drei Funktionen, die durch ihre Struktur und ihre Basensequenz bestimmt werden (Abb. 101).

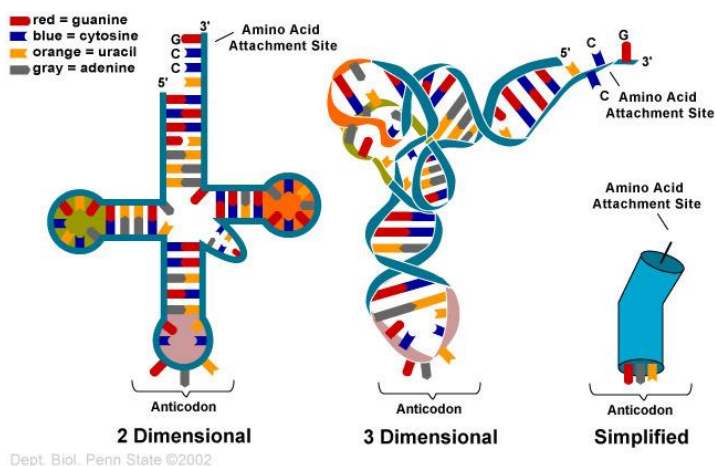


Abbildung 101 tRNA

1. tRNAs binden spezifische Aminosäuren. Jede tRNA bindet an ein Enzym, das spezifisch eine der 20 Aminosäuren an sie koppelt. Die kovalente Bindung mit der Aminosäure erfolgt am 3'-Ende der tRNA. Wenn die tRNA eine Aminosäure trägt, bezeichnet man sie als „beladen“. Die dafür zuständigen Enzyme werden als Aminoacyl-tRNA-Synthetasen bezeichnet. Jedes dieser Enzyme ist für eine Aminosäure und die zugehörige tRNA spezifisch. Bei der Reaktion wird ATP verbraucht, sodass zwischen einer tRNA und der Aminosäure eine sehr energiereiche Bindung entsteht. Die Energie dieser Bindung dient später dazu, die Peptidbindung zwischen der Aminosäure und der wachsenden Polypeptidkette zu bilden.
2. tRNAs binden an mRNA. Etwa in der Mitte der Polynucleotidkette der tRNA befindet sich ein Triplet, das man als Anticodon bezeichnet. Es ist komplementär zum mRNA-Codon für die zugehörige Aminosäure, die die betreffende tRNA trägt. Wie die beiden Stränge der DNA binden auch das Codon und das Anticodon über Wasserstoffbrücken aneinander.

3. tRNAs treten mit dem Ribosom in Wechselwirkung. Das Ribosom enthält an seiner Oberfläche mehrere Stellen, die genau zur dreidimensionalen Struktur des tRNA-Moleküls passen. Die komplexe Struktur des Ribosoms ermöglicht, dass es die mRNA und die beladenen tRNAs in den richtigen Positionen festhalten kann, sodass eine Polypeptidkette effizient zusammengebaut wird. Ein Ribosom kann jede beliebige mRNA und alle Typen von beladenen tRNAs verwenden und eignet sich daher zur Herstellung von Polypeptiden aller Art. Ribosomen können immer wieder verwendet werden, und in den meisten Zelltypen gibt es Tausende davon.

Ein Ribosom besteht aus zwei Untereinheiten, einer großen und einer kleinen. Diese beiden Untereinheiten und mehrere Dutzend weitere Moleküle interagieren nichtkovalent (Abb. 102).

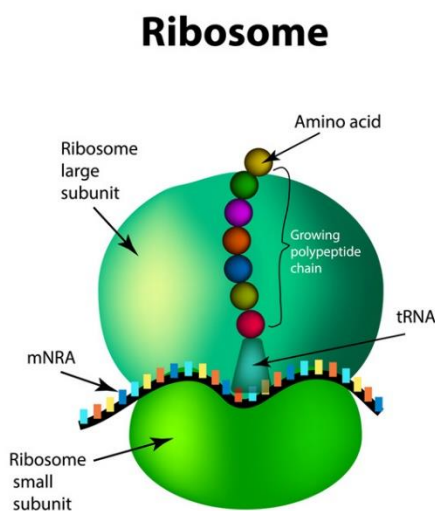


Abbildung 102 Ribosom

Bei den Eukaryoten setzt sich die große Untereinheit aus drei verschiedenen Molekülen ribosomaler RNA (rRNA) und 49 verschiedenen Proteinmolekülen zusammen, die in einer ganz bestimmten Weise angeordnet sind. Die kleine Untereinheit umfasst ein rRNA-Molekül und 33 verschiedene Proteinmoleküle.

Die Ribosomen der Prokaryoten sind etwas kleiner als die der Eukaryoten, und ihre Proteine und rRNAs unterscheiden sich etwas von denen der Eukaryoten.

Auf der großen ribosomalen Untereinheit befinden sich drei Stellen, an die eine tRNA binden kann. Sie werden mit A, P und E bezeichnet (Abb. 103). Die mRNA und das Ribosom bewegen sich relativ zueinander, und währenddessen durchwandert die beladene tRNA diese drei Stellen in einer bestimmten Reihenfolge:

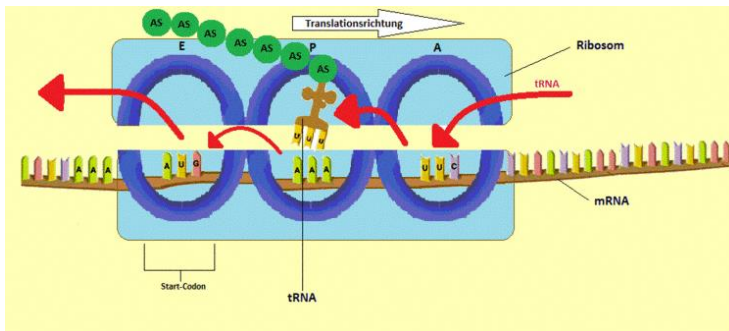


Abbildung 103 A-, P-, E-Stelle des Ribosoms

1. An der A-(Aminoacyl-)Stelle bindet das Anticodon der beladenen tRNA an das mRNA-Codon und bringt so die richtige Aminosäure an die wachsende Polypeptidkette.
2. An der P-(Peptidyl-)Stelle hängt die tRNA ihre Aminosäure an die Polypeptidkette.
3. An der E-(Exit-)Stelle befindet sich die tRNA, nachdem sie ihre Aminosäure abgegeben hat, bevor sie vom Ribosom freigesetzt wird und in das Cytosol zurückkehrt. Dort nimmt sie dann eine neue Aminosäure auf und der Vorgang beginnt von vorne.

Das Ribosom besitzt eine Genauigkeitsfunktion, die sicherstellt, dass die Wechselwirkungen zwischen tRNA und mRNA korrekt erfolgen. Das heißt, dass eine beladene tRNA äußerst zuverlässig mit dem richtigen Anticodon an das passende Codon bindet. Wenn es zu einer ordnungsgemäßen Bindung kommt, bilden sich zwischen den Basenpaaren Wasserstoffbrücken. Die rRNA der kleinen Untereinheit ist daran beteiligt, das Zusammenpassen der drei Basen zu überprüfen. Wenn sich nicht zwischen allen drei Basenpaaren Wasserstoffbrücken gebildet haben, muss die tRNA für dieses mRNA-Codon falsch sein und wird aus dem Ribosom entfernt.

Die Translation erfolgt in drei Schritten, die man – wie bei der Transkription – als Initiation, Elongation und Termination bezeichnet (Abb. 104).

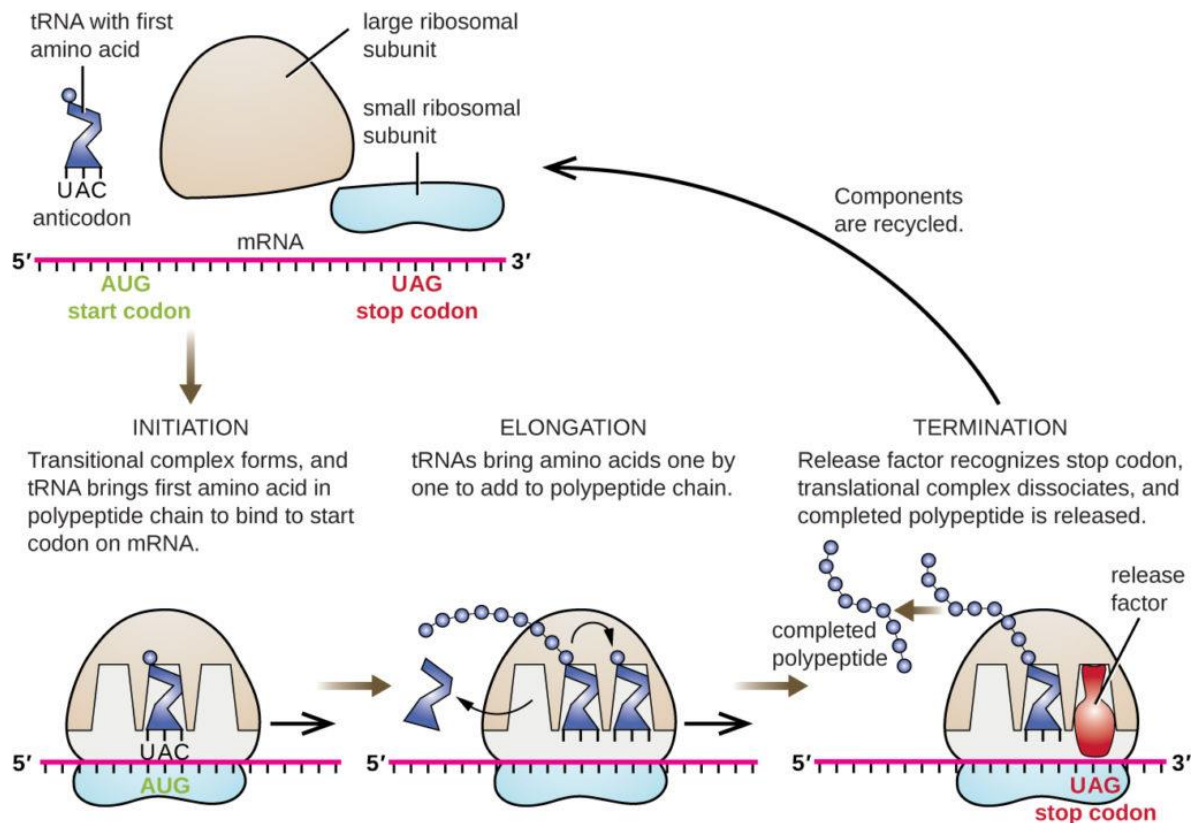


Abbildung 104 Translation

Initiation:

Die Translation einer mRNA beginnt mit der Bildung eines Initiationskomplexes, der aus einer beladenen tRNA und einer kleinen ribosomalen Untereinheit besteht, die beide an die mRNA gebunden sind. Die kleine Untereinheit bindet bei Eukaryoten an die 5'-Cap-Gruppe der mRNA. Bei Prokaryoten bindet die rRNA der kleinen Untereinheit zuerst an eine komplementäre Ribosomenbindungsstelle (AGGAGG,) auf der mRNA. Diese Sequenz liegt weniger als zehn Nucleotide stromaufwärts des eigentlichen Startcodons und ordnet das Startcodon so an, dass es in der Nähe der P-Stelle der großen Untereinheit liegt.

Nach der Bindung wandert die kleine Untereinheit die mRNA entlang, bis sie auf das Startcodon trifft. Das Startcodon hat die Sequenz AUG und die dazu passende tRNA hat die Aminosäure Methionin im Schlepptau. Die erste Aminosäure eines sich bildenden Polypeptids ist also immer ein Methionin. Nach der Bindung der mit Methionin beladenen tRNA an die mRNA tritt die große ribosomale Untereinheit hinzu. Die mit Methionin beladene tRNA befindet sich nun in der P-Stelle des Ribosoms, und in der A-Stelle liegt das zweite mRNA-Codon.

Elongation:

Eine beladene tRNA, deren Anticodon zum zweiten Codon der mRNA komplementär ist, tritt nun in die offene A-Stelle der großen ribosomalen Untereinheit ein. Dann katalysiert die große Untereinheit zwei Reaktionen:

1. Die Bindung zwischen der tRNA und ihrer Aminosäure in der P-Stelle wird gelöst.
2. Zwischen dieser Aminosäure und derjenigen, die an die in der A-Stelle befindliche tRNA gebunden ist, wird eine Peptidbindung gebildet.

Da die große ribosomale Untereinheit also einen Peptidrest auf eine Aminosäure überträgt (transferiert), spricht man hier von einer Peptidyltransferaseaktivität. Katalysator dieser Reaktion ist die rRNA, weil Experimente, bei denen die Proteinkomponenten der Ribosomen entfernt wurden, nicht jedoch die rRNA, zeigten, dass die Peptidbindung trotzdem weiter stattfand. Die rRNA fungiert also als Ribozym. Nachdem die Peptidbindung zwischen Aminosäure eins und zwei an der A-Stelle geknüpft wurde, bewegt sich das Ribosom entlang der mRNA in 5'-3'-Richtung um ein Codon weiter. Die tRNA mit dem Peptidylrest aus zwei Aminosäuren gelangt zur P-Stelle. Die freie tRNA, die mit der ersten Aminosäure beladen war, befindet sich nun an der E-Stelle und verlässt das Ribosom. Sie kann im Cytosol erneut mit Aminosäuren beladen werden. Die A-Stelle ist frei, sodass die nächste tRNA an die mRNA binden kann. Der Vorgang der Elongation setzt sich fort und die Polypeptidkette wächst, während sich diese Schritte wiederholen.

Termination:

Der Elongationszyklus endet und die Translation wird terminiert, sobald ein Stoppcodon – UAA, UAG oder UGA – in die A-Stelle eintritt. Stoppcodons codieren keine Aminosäuren und binden auch keine tRNA. Sie binden vielmehr einen Freisetzungsfaktor (release-Faktor), der die Bindung zwischen dem entstandenen Polypeptid und der tRNA an der P-Stelle hydrolysiert. Das neu gebildete Polypeptid löst sich nun vom Ribosom. Sein C-Terminus ist die letzte Aminosäure, die an die Kette angehängt wurde. Der N-Terminus ist (zumindest zunächst) ein Methionin (als Folge des AUG-Startcodons).

Kapitel 11: Genregulation und Epigenetik

In den vorherigen Beiträgen haben wir den Bau der DNA kennengelernt und ihre Rolle bei der Zellteilung und der Herstellung von Proteinen. Die menschliche Zelle hat aber über 20.000 Gene und nicht jede Zelle unseres Körpers ist gleich. Muskelzellen brauchen andere Proteine als Nervenzellen usw. Es wäre ziemlich ineffizient, wenn alle Gene gleichzeitig in jeder Zelle exprimiert werden. Es muss also ein System vorhanden sein, dass Gene nur in den Zellen exprimiert werden, wo sie benötigt werden und auch zum richtigen Zeitpunkt. Dies wird mittels der Genregulation gesteuert. Je früher die Zelle in den Prozess der Proteinsynthese eingreift, umso weniger Energie und Aminosäuren verschwendet sie.

Die Genexpression beginnt am Promotor, wo die RNA-Polymerase bindet und dann mit der Transkription beginnt (Abb. 105). Dabei ist festzuhalten, dass in einer beliebigen Zelle zu einem beliebigen Zeitpunkt niemals alle Promotoren aktiv sind. Das deutet darauf hin, dass die Transkription von Genen selektiv erfolgen muss.

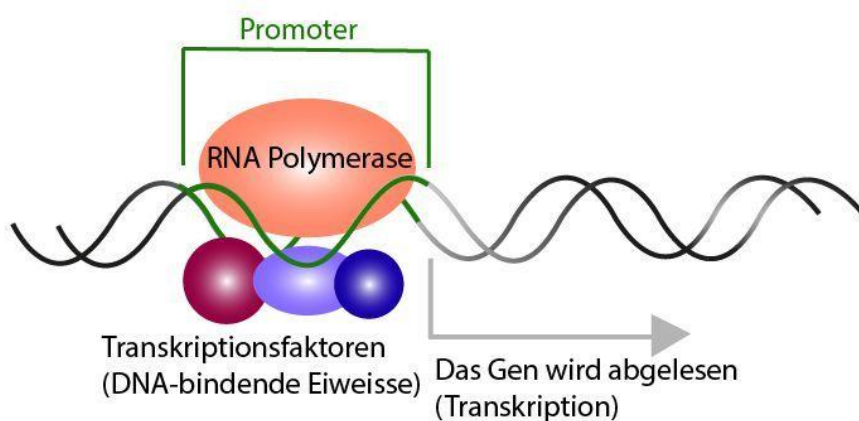


Abbildung 105 RNA-Polymerase bindet an den Promotor

Bei der „Entscheidung“, welche Gene zu aktivieren sind, wirken zwei Arten von regulatorischen Proteinen mit, die an die DNA binden: Repressorproteine und Aktivatorproteine. Beide binden an den Promotor und regulieren dadurch das Gen.

Bei der negativen Regulation verhindert die Bindung des Repressorproteins an die DNA die Transkription.

Bei der positiven Regulation bindet ein Aktivatorprotein an die DNA und stimuliert dadurch die Transkription.

Genregulation bei Prokaryoten

Die ersten Erkenntnisse der Genregulation wurden bei Bakterien untersucht, vor allem bei *Escherichia coli*. Als normaler Bewohner des menschlichen Darms muss es sich an schnelle Veränderungen seiner chemischen Umgebung anpassen.

Der Wirt kann die Bakterien im stündlichen Wechsel mit andersartigen Nährstoffen konfrontieren (z. B. nacheinander mit Glucose und Lactose). Solche Veränderungen der Nährstoffe stellen spezifische Anforderungen an den Stoffwechsel des Bakteriums. Glucose (Traubenzucker) ist die von ihm bevorzugte Energiequelle; dieser Zucker lässt sich am einfachsten umsetzen. Fehlt Glucose als Nährstoff, kann es Lactose verdauen. Damit die Lactose von *E. coli* aufgenommen und verarbeitet werden kann, sind drei Enzyme erforderlich. Das Bakterium verschwendet keine Energie diese Proteine zu bilden, wenn im Wachstumsmedium keine Laktose vorhanden ist. Diese Proteine werden nur hergestellt, wenn Laktose für den Stoffwechsel verfügbar ist. Wie schafft es das Bakterium die Gene für die Lactose-abbauenden Enzyme zu aktivieren?

Dies wird mit dem sog. lac-Operon Modell erklärt (Abb. 106). Dieses besteht aus dem Promotor, als der Stelle an der die RNA-Polymerase bindet. Der zweite Bestandteil ist der Operator. Dies ist quasi gesehen der „Schalter“ für die Genaktivität. Am Operatorsequenzabschnitt findet sich eine Bindestelle für das Regulatorprotein. Die Transkription durch die RNA-Polymerase kann durch diese Blockade verhindert werden. der dritte Bestandteil sind die Strukturgene, im Falle des lac-Operons die drei Enzyme zum Abbau der Lactose. Die Strukturgene werden als lacZ, lac Y und lacA bezeichnet. Für die eigentliche Regulation verantwortlich ist das Regulator- oder Repressorprotein R. Das Gen für den Repressor (lacI) liegt außerhalb des lac-Operons. Wenn der Repressor am Operator des lac-Operons gebunden hat, wird die Transkription des Operons blockiert.

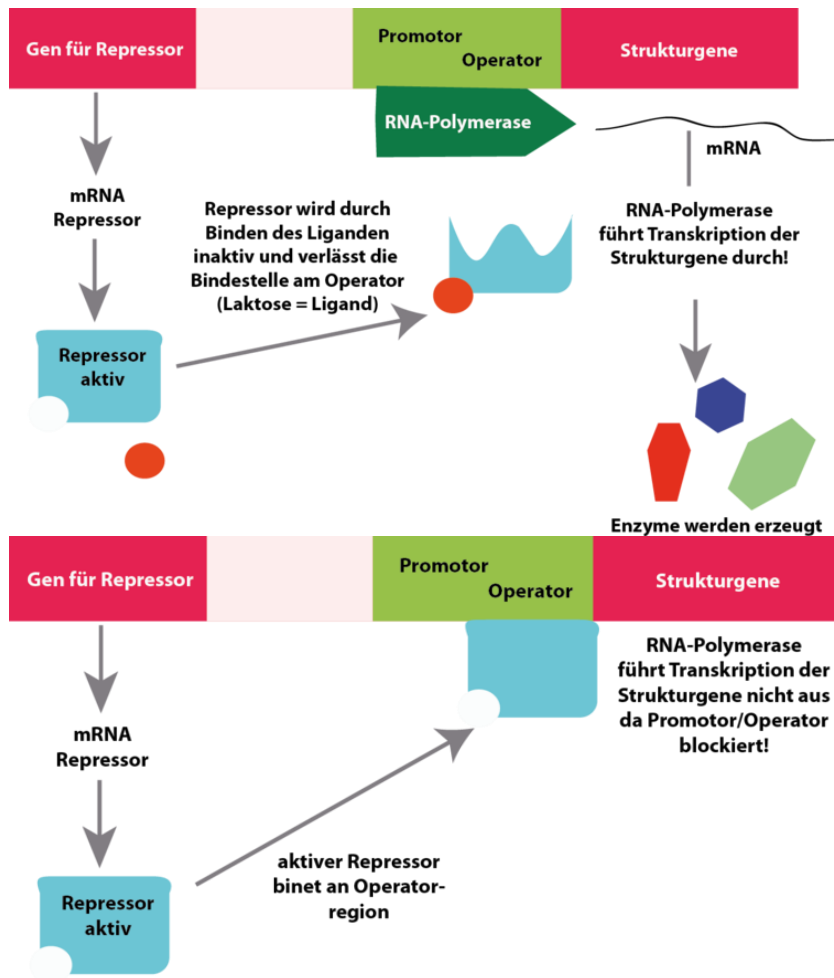


Abbildung 106 Lac-Operon. Oben: Laktose bindet an das Repressorprotein und aktiviert das lac-Operon; die Gene für die Enzyme zur Laktose-Spaltung werden aktiviert.

Unten: Bei Abwesenheit von Laktose kann das Repressorprotein an den Operator binden und blockiert das Ablesen der Gene am Lac-Operon.

Das Lac-Operon wird durch das Substrat Lactose aktiviert. Man spricht von Substratinduktion.

Das Repressorprotein enthält zwei Bindungsstellen: eine für den Operator und eine für den Induktor. Das äußere Signal aus der Umgebung, welches das lac-Operon von *E. coli* induziert ist Lactose. Der eigentliche Induktor ist jedoch Allolactose, ein Molekül, das aus Lactose gebildet wird, sobald Lactose in die Bakterienzelle gelangt. Ist der Induktor nicht vorhanden, bindet das Repressorprotein an den Operator. Dadurch kann die RNA-Polymerase nicht an den Promotor binden und das Operon wird nicht transkribiert. Sobald der Induktor vorhanden ist, bindet er an den Repressor und verändert dessen Konformation. Diese Änderung der dreidimensionalen Struktur verhindert, dass der Repressor an den Operator bindet. Dadurch kann die RNA-Polymerase an den Promotor binden und mit der Transkription der Strukturgene des lac-Operons beginnen. Also: Das Vorhandensein von Lactose unterdrückt den Repressor, der nicht an den Operator binden kann und ermöglicht das Ablesen der Strukturgene. Man bezeichnet dies als Substratinduktion.

Neben der Substratinduktion gibt es auch die Endproduktrepression. Hier bindet ein Stoffwechselprodukt an ein regulatorisches Protein, das dann an den Operator binden kann und die Transkription blockiert.

Ein Beispiel dafür ist das Operon, dessen Strukturgene die Synthese der Aminosäure Tryptophan katalysieren (Abb. 107). Wenn Tryptophan in der Zelle in ausreichender Konzentration vorkommt, ist es ökonomisch, die Produktion der Enzyme für die Tryptophansynthese abzuschalten. Dafür verwendet die Zelle einen Repressor, der an einen Operator im trp-Operon bindet.

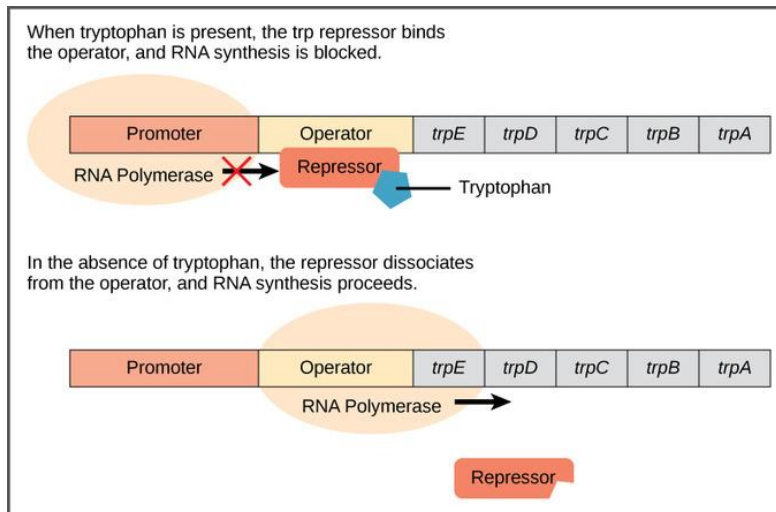


Abbildung 107 trp-Operon

Aber der Repressor des trp-Operons ist normalerweise nicht an den Operator gebunden. Er bindet nur, wenn sich seine Struktur durch die Bindung mit Tryptophan verändert.

Das lac-Operon und trp-Operon sind Beispiele für die Funktion von Repressorproteinen, die also an den Operator binden und die Transkription der Strukturgene blockieren. Es gibt aber natürlich auch Aktivatorproteine, die die Transkription stimulieren.

Wie wir erfahren haben, kann der lac-Repressor nicht an den lac-Operator binden und so die Transkription zu blockieren, wenn Lactose vorhanden ist. Glucose wird jedoch von der Zelle als Energiequelle bevorzugt, sodass das lac-Operon trotz eines hohen Lactosespiegels nicht wirkungsvoll transkribiert wird, wenn zugleich der Glucosespiegel noch relativ hoch sind. Das liegt daran, dass für die effiziente Transkription des lac-Operons zusätzlich die Bindung eines Aktivatorproteins an den Operator notwendig ist.

Bei einer geringen Glucosekonzentration in der Zelle wird ein Signalweg angeschaltet, der dazu führt, dass die Konzentration von cAMP (ein sekundärer Botenstoff) zunimmt. Zyklisches AMP bindet an das Aktivatorprotein CRP (cAMP response protein). Dadurch ändert sich die Konformation von CRP, sodass das Protein an den Promotor des lac-Operons binden kann. CRP ist ein Transkriptionsaktivator, da seine Bindung dazu führt, dass die RNA-Polymerase leichter an den Promotor andockt und sich dadurch die Transkriptionsrate erhöht. In Gegenwart einer ausreichenden Menge

Glucose ist die cAMP-Konzentration gering, sodass CRP nicht an den Promotor bindet, und die Transkription des lac-Operons ist weniger effizient.

Genregulation bei Eukaryoten

Damit nicht genug. Untersuchungen der Genregulation bei Bakterien haben deren Grundlagen entschlüsselt. Bei Eukaryoten und insbesondere bei mehrzelligen Lebewesen ist die Genregulation um einiges komplexer. Zum einen ist das Genom komplexer und zum anderen gibt es verschiedene Zelltypen mit verschiedenen Aufgaben und damit Genexpressionen. Bei Eukaryoten spielen die Transkriptionsfaktoren eine sehr wichtige Rolle.

Eukaryoten nutzen für die Transkription natürlich ebenfalls eine RNA-Polymerase. Es gibt jedoch in eukaryotischen Zellen hiervon drei verschiedene: Die RNA-Polymerase I, II und III. Typ I transkribiert die rRNA, Typ III die tRNA und weitere RNAs und Typ II die mRNA. Die RNA-Polymerase II kann jedoch nicht an DNA binden und selbstständig mit der Transkription beginnen. In Eukaryoten muss eine Reihe von Proteinen, sogenannte Transkriptionsfaktoren, an einen Promotor binden, bevor die RNA-Polymerase II die Transkription initiieren kann (Abb. 108).

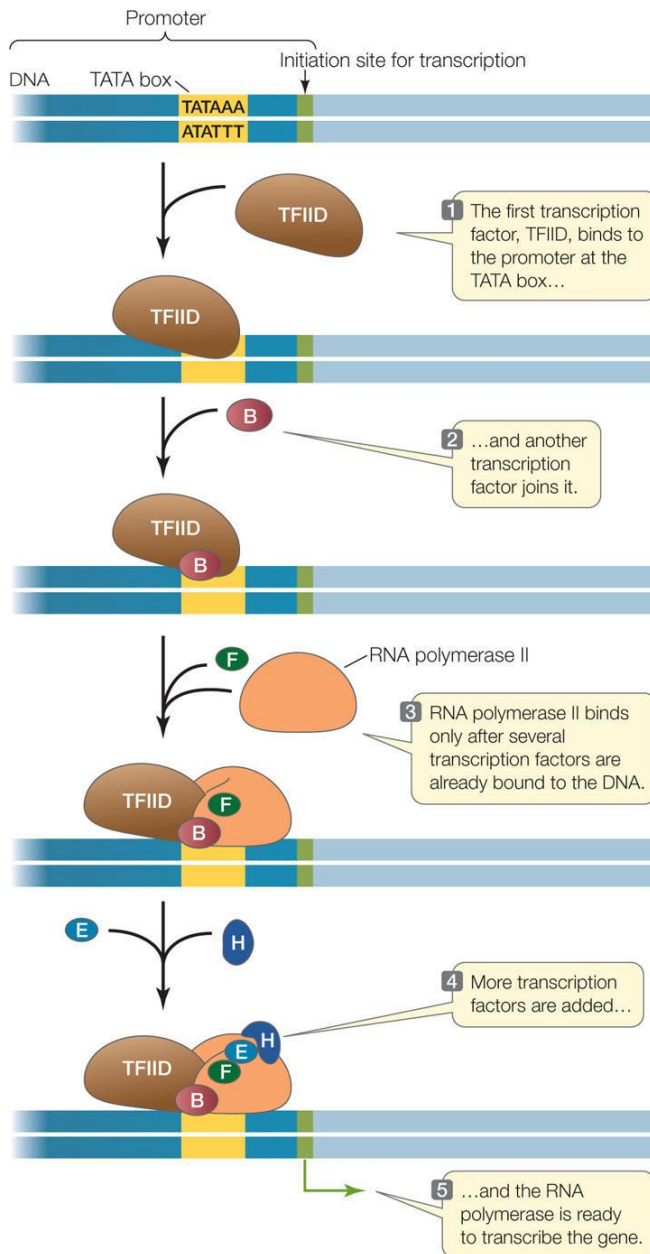


Abbildung 108 Genregulation bei Eukaryoten

Der erste Transkriptionsfaktor, der bindet, wurde TFIID genannt. TFIID bindet an eine Region des Promotors, die als TATA-Box bezeichnet wird und die Sequenz TATA enthält. Die TATA-Box ist diejenige Stelle, an der die DNA anfängt zu denaturieren, sodass der Matrizenstrang zugänglich wird. Nach der TFIID-Bindung binden eine Reihe anderer Transkriptionsfaktoren und die RNA-Polymerase II einen Initiationskomplex.

Regulatorische DNA-Sequenzen wie die TATA-Box kommen in den Promotoren von zahlreichen eukaryotischen Genen vor; sie werden von Transkriptionsfaktoren erkannt, die in allen Zellen eines Organismus vorhanden sind. Andere regulatorische Sequenzen kommen nur in wenigen Genen vor und werden von spezifischen Transkriptionsfaktoren erkannt. Diese Faktoren kommen möglicherweise nur in bestimmten Zelltypen oder in bestimmten Phasen des Zellzyklus vor, oder sie werden

nur als Reaktion auf Signale in der Zelle oder aus der Umgebung über Signalwege aktiviert.

Im Promotor kann sich z. B. eine DNA-Region befinden, die an Regulatorproteine binden kann. Ein Regulatorprotein bindet an die Regulatorregion und an den Transkriptionsinitiations-komplex und aktiviert ihn dadurch.

Einige regulatorische DNA-Sequenzen sind positive Elemente, die man als Enhancer bezeichnet: An sie binden Transkriptionsfaktoren, die entweder die Transkription aktivieren oder die Transkriptionsrate erhöhen. Andere regulatorische Elemente sind Silencer: An sie binden Faktoren, die die Transkription unterdrücken (Abb. 109).

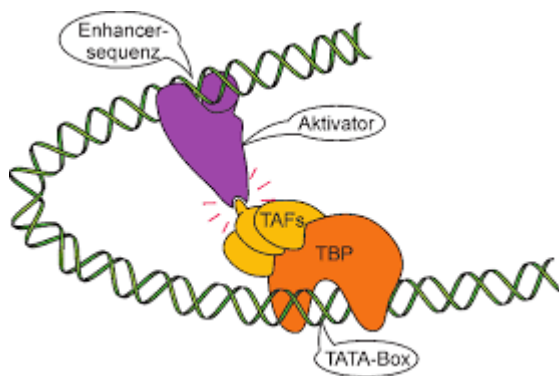


Abbildung 109 Enhancer

Die meisten regulatorischen Elemente, die für die korrekte Expression eines Gens erforderlich sind, liegen innerhalb weniger Hundert Basenpaare stromaufwärts des Transkriptionsstartpunkts.

So enthält beispielsweise der Promotor des Albumingens der Maus die gesamte Information, die für die spezifische Expression des Gens in Leberzellen notwendig ist, innerhalb von 170 bp stromaufwärts des Transkriptionsstartpunkts.

Einige regulatorische Elemente können aber durchaus Tausende von Basenpaaren entfernt liegen und die Expression von mehreren nahe beieinander liegenden Genen beeinflussen. Wenn Transkriptionsfaktoren an diese Elemente binden, treten sie mit dem RNA-Polymerase-Komplex in Wechselwirkung und verursachen eine Krümmung der DNA. Die Kombination aus der Bindung von Transkriptionsfaktoren an ein Gen bestimmt die Transkriptionsrate. So enthält beispielsweise ein unreifes rotes Blutkörperchen (Erythrocyt) im Knochenmark große Mengen an β -Globin. In diesen Zellen wirken mindestens 13 verschiedene Transkriptionsfaktoren bei der Regulation der Transkription des β -Globin-Gens mit. Nicht alle diese Faktoren sind in anderen Zellen aktiv oder überhaupt vorhanden, beispielsweise in unreifen weißen Blutkörperchen (Leukocyten), die vom selben Knochenmark produziert werden. Das führt dazu, dass das β -Globin-Gen in diesen Zellen nicht exprimiert wird. Es sind zwar in allen Zellen die gleichen Gene vorhanden, aber der Werdegang einer Zelle wird dadurch bestimmt, welche Gene sie wann exprimiert.

Wie aber erkennen Transkriptionsfaktoren spezifische DNA-Sequenzen? Transkriptionsfaktoren haben spezifische Domänen, die an die DNA binden können. In diesen Proteindomänen gibt es mehrere allgemeine Strukturmotive, die an DNA

binden. Diese bestehen aus verschiedenen Kombinationen von Sekundärstrukturen; sie können auch besondere Komponenten wie Zink enthalten. Eines der allgemeinen Struktur motive in DNA-bindenden Domänen ist das Helix-Turn-Helix-Motiv. Die in der großen DNA-Furche liegende „Erkennungshelix“ tritt mit den DNA-Basen sequenzspezifisch in Wechselwirkung. Die im rechten Winkel dazu nach außen zeigende Helix interagiert mit dem Zucker-Phosphat-Rückgrat, wodurch sichergestellt ist, dass die innere Helix in der richtigen Orientierung mit den Basen in Kontakt tritt. Wie kann ein Protein in der DNA eine Sequenz erkennen? Wie wir bereits wissen, bilden die komplementären Basen der DNA Wasserstoffbrücken untereinander. Sie können aber auch zusätzliche Wasserstoffbrücken mit Proteinen bilden. Auf diese Weise kann die intakte DNA-Doppelhelix von einem Proteinmotiv erkannt werden, dessen Struktur in die große oder kleine Furche hineinpasst, Aminosäuren enthält, die in das Innere der Doppelhelix hineinragen können, sowie Aminosäuren enthält, die mit den innen liegenden Basen Wasserstoffbrücken bilden können (Abb. 110).

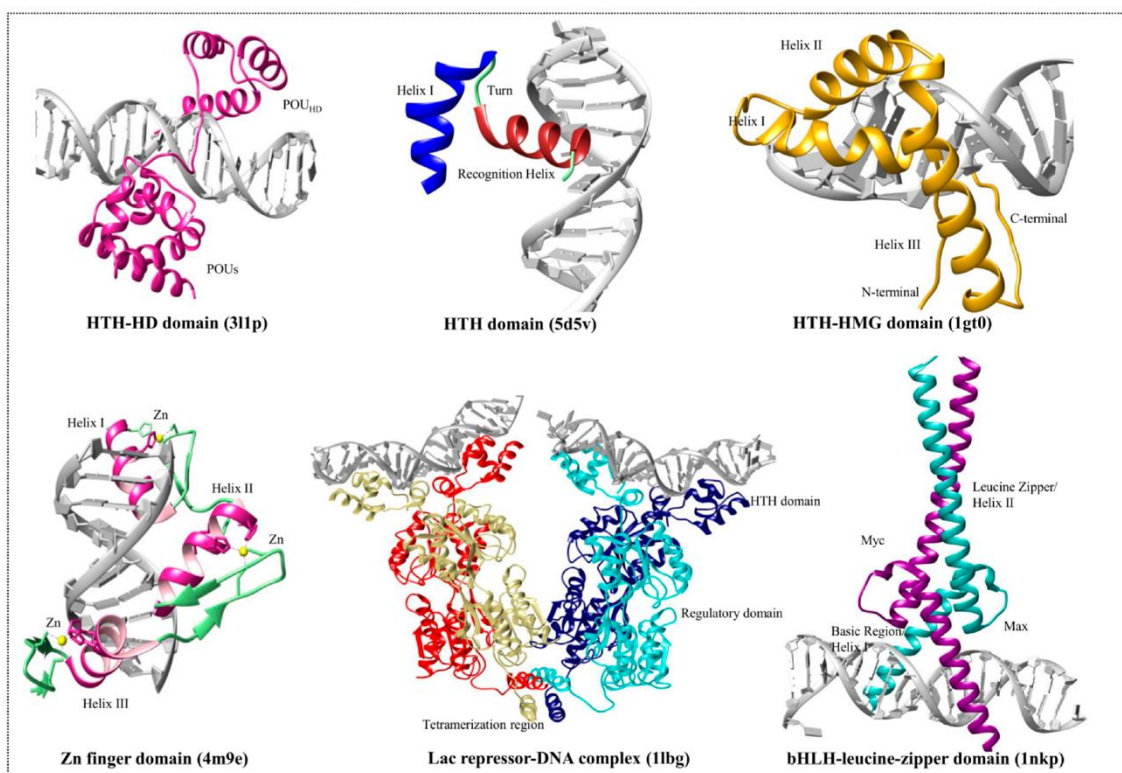


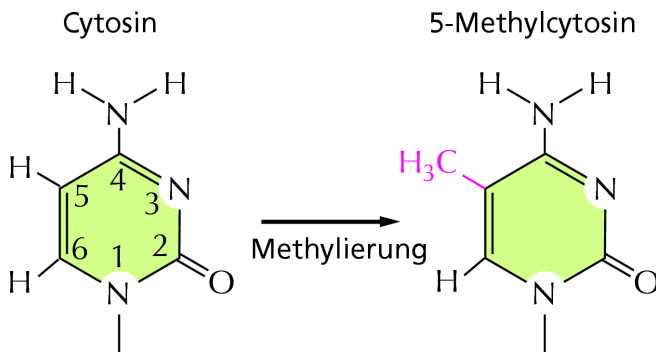
Abbildung 110 Transkriptionsfaktoren bin an DNA

Epigenetik

Die Genregulation hört aber bei Transkriptionsfaktoren keineswegs auf. Immer wichtiger werden die Erkenntnisse der Epigenetik. In der Mitte des 20. Jahrhunderts formte der Entwicklungsbiologe Conrad Waddington den Begriff „Epigenetik“. Diese definierte er als „den Zweig der Biologie, der sich mit den kausalen Wechselwirkungen zwischen Genen und ihren Produkten beschäftigt, die den Phänotyp hervorbringen“. Heute definiert man die Epigenetik spezifischer und meint damit die Veränderung der Genexpression, die ohne eine Veränderung der DNA-Sequenz erfolgt. Diese

Veränderungen sind reversibel, manchmal aber stabil und vererbbar. Die zwei bekanntesten Formen der Epigenetik sind DNA-Methylierung und Histon-Modifikation.

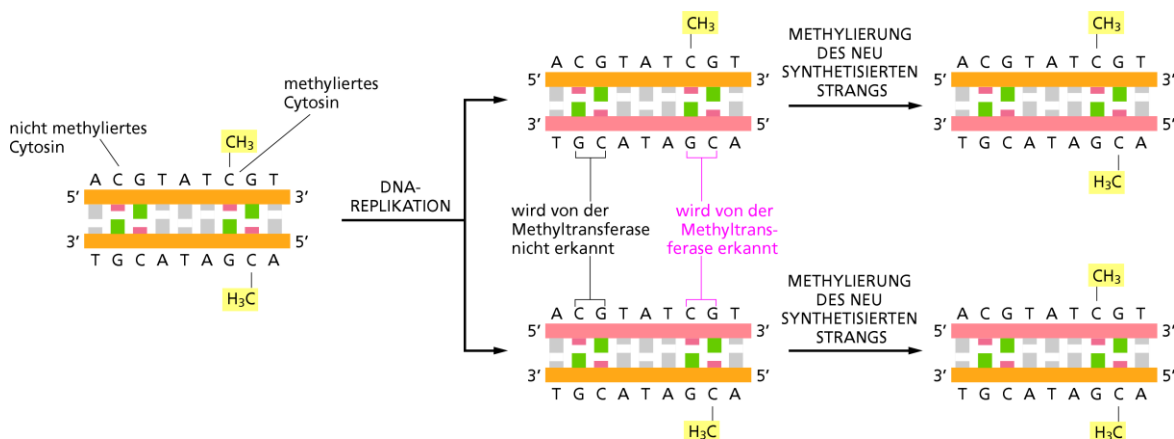
Abhängig vom Organismus werden in der DNA 1–5% der Cytosine durch Anhängen einer Methylgruppe ($-\text{CH}_3$) zu 5-Methylcytosin modifiziert (Abb. 111). Diese kovalente Bindung wird von dem Enzym DNA-Methyltransferase katalysiert und tritt bei Säugern im Allgemeinen bei solchen Cytosinen (C) auf, die in der Sequenz stromaufwärts neben einem Guanin (G) liegen. Diese DNA-Regionen bezeichnet man als CpG-Inseln; sie kommen vor allem in Promotoren häufig vor.



© 2012 Wiley-VCH, Weinheim
Alberts - Lehrbuch der Molekularen Zellbiologie
ISBN: 978-3-527-32824-6 Fig. 08-021

Abbildung 111 Methylierung von Cytosin

Diese kovalente Veränderung der DNA ist während der Mitose erblich: Wenn die DNA repliziert wird, katalysiert eine besondere Methyltransferase die Bildung von 5-Methylcytosin im neuen DNA-Strang. Dieses Enzym führt also „Wartungsarbeiten“ durch, indem es dort, wo ein DNA-Strang Methylcytosin enthält, nach dessen Replikation die Methylierung kopiert (Abb. 112). Man spricht hier deshalb von maintenance-Enzymen („Wartungsenzymen“). Das Muster der Cytosinmethylierung kann sich auch ändern, da die Methylierung reversibel ist. Ein drittes Enzym, die Demethylase, katalysiert das Entfernen der Methylgruppe von Cytosin.



© 2012 Wiley-VCH, Weinheim
Alberts - Lehrbuch der Molekularen Zellbiologie
ISBN: 978-3-527-32824-6 Fig. 08-022

Abbildung 112 Methylierung der DNA bei der Replikation

Welche Auswirkungen hat die DNA-Methylierung? Während der Replikation und der Transkription verhält sich 5-Methylcytosin wie das nicht-methylierte Cytosin. Es bildet

Basenpaare mit Guanin. Aber zusätzliche Methylgruppen in einem Promotor rekrutieren weitere Proteine, die an methylierte DNA binden. Diese Proteine sind generell an der Repression der Gentranskription beteiligt. Stark methylierte Gene sind in der Tendenz inaktiver. Diese Art der genetischen Regulation ist epigenetisch, da Genexpressionsmuster ohne Veränderung der DNA-Sequenz beeinflusst werden.

Die DNA-Methylierung ist für die Entwicklung vom Ei zum Embryo von Bedeutung. Wenn beispielsweise ein Spermienzellkern eines Säugers in eine Eizelle eindringt, werden zuerst im männlichen, danach im weiblichen Genom zahlreiche Gene demethyliert. Viele Gene, die also normalerweise inaktiv sind, werden während der frühen Entwicklungsphase exprimiert.

Während sich der Embryo entwickelt und sich seine Zellen spezialisieren, werden Gene, deren Produkte für bestimmte Zelltypen nicht notwendig sind, methyliert. Diese methylierten Gene sind stillgelegt, ihre Transkription ist reprimiert. Ungewöhnliche oder anormale Ereignisse können jedoch stillgelegte Gene wieder einschalten.

Ein weiterer Mechanismus für eine epigenetische Genregulation ist die Veränderung der Chromatinstruktur (Chromatin-Remodeling). Wie wir erfahren haben, ist die DNA mit Histonproteinen zu Nucleosomen verpackt, sodass die DNA für die RNA-Polymerase und den übrigen Transkriptionsapparat unzugänglich ist. Jedes Histonprotein enthält an seinem N-Terminus einen „Schwanz“ aus etwa 20 Aminosäuren, der aus der kompakten Struktur herausragt und an bestimmten Positionen positiv geladene Aminosäuren besitzt (vor allem Lysin). Normalerweise besteht zwischen positiv geladenen Histonproteinen und der DNA, die aufgrund ihrer Phosphatgruppen negativ geladen ist, eine starke ionische Anziehungskraft.

Jedoch können Enzyme mit der Bezeichnung Acetyltransferasen an diesen positiv geladenen Aminosäuren Acetylgruppen befestigen, sodass sich die Ladung ändert (Histonacetylierung). Wenn die positive Ladung an den Histonschwänzen verringert wird, nimmt auch die Affinität der Histone für die DNA ab, sodass sich die kompakten Nucleosomen öffnen.

Weitere Chromatin-Remodeling-Proteine können an die gelockerten Histon-DNA-Komplexe binden. So wird die DNA für die Genexpression geöffnet. Histon-Acetyltransferasen können also die Transkription aktivieren (Abb. 113).

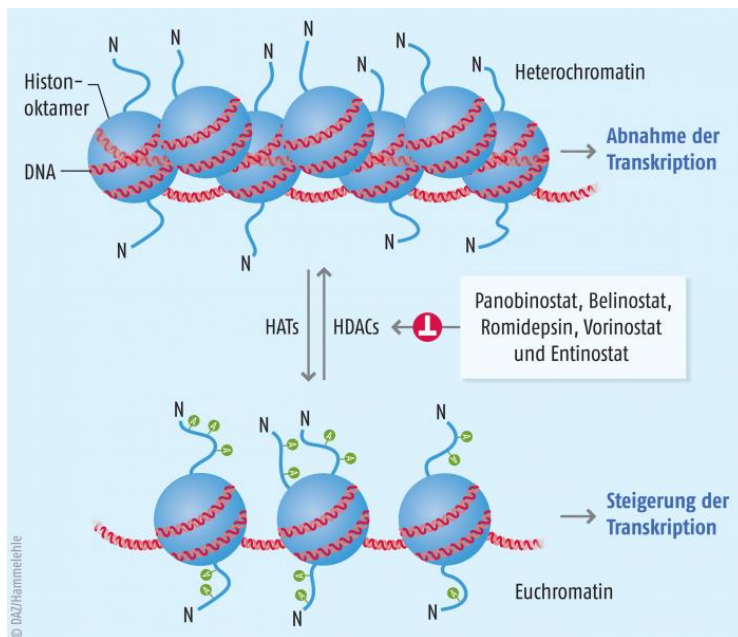


Abbildung 113 Histon-Modifikation

Ein weiterer Typ von Remodelingproteinen sind die Histon-Deacetylasen. Sie können Acetylgruppen von Histonen entfernen und reprimieren so die Transkription.

Acetylierung ist nicht die einzige Form der Modifikation von Histonen, durch den die Genaktivierung und die Genrepression beeinflusst werden. So hängt beispielsweise die Histonmethylierung (die Sie nicht mit der Methylierung der DNA verwechseln sollten) mit der Inaktivierung von Genen zusammen. Auch die Histonphosphorylierung beeinflusst die Genexpression.

Die spezifische Wirkung hängt davon ab, welche Aminosäure modifiziert wird. Alle diese Effekte sind reversibel, sodass die Aktivität eines eukaryotischen Gens durch sehr komplexe Muster der Histonmodifikation bestimmt sein kann.

Epigenetische Veränderungen werden z. B. durch Umweltfaktoren beeinflusst. Z. B. besteht ein Bienenstaat aus verschiedenen Kasten: unfruchtbare Arbeiterbienen, fruchtbare Königinnen, etc. Was bestimmt aber, ob aus einer Bienenlarve eine Königin oder Arbeiterin wird? Sicherlich nicht die Erbfolge der DNA, denn weibliche Honigbienen besitzen alle eine sehr ähnliche genetische Ausstattung. Während des Larvenstadiums füttern die Arbeiterinnen jedoch bestimmte Weibchen im Stock mit Gelée royale (eine proteinreiche Substanz), wodurch sich bei zahlreichen Genen die Expression ändert. Die jungen Königinnen werden viel größer als die Arbeiterinnen und wird geschlechtsreif (Abb. 114).



Epigenetische Veränderungen sind zwar reversibel, aber viele dieser Veränderungen, etwa DNA-Methylierung und Histonacetylierung, können die Genexpressionsmuster in einer Zelle dauerhaft verändern.

Alternatives Spleißen und RNA-Interferenz

144

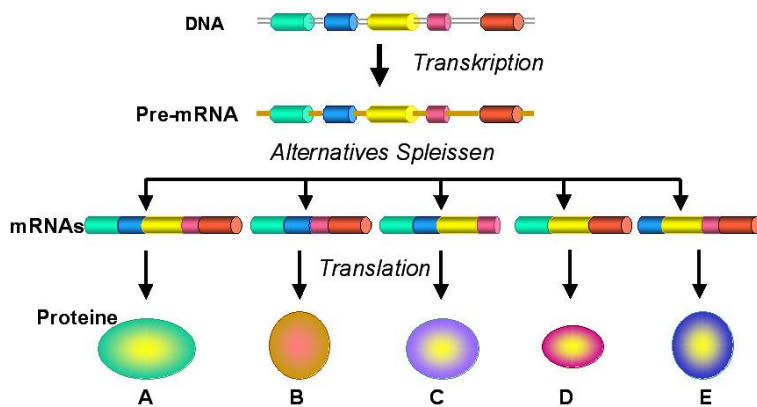
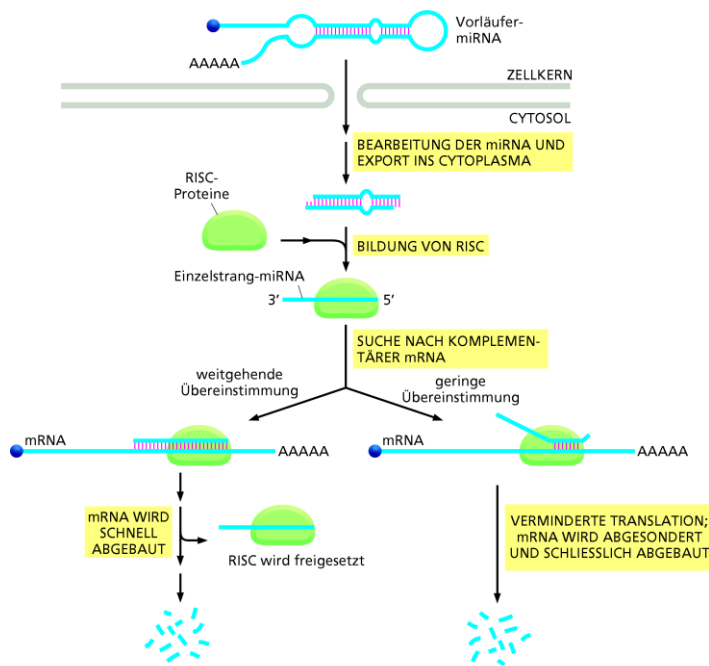


Abbildung 115 Alternatives Spleißen

Nehmen wir mal an ein Gen besteht aus 5 Exons und 4 dazwischenliegenden Introns. Beim Spleißen würden die 4 Introns ausgeschnitten werden, sodass die 5 Exons übrigbleiben. Beim normalen Spleißen würden die Exons in der „richtigen“ Reihenfolge 1-5 hintereinander zusammengefügt werden. Beim alternativen Spleißen können die Exons jedoch unterschiedlich kombiniert werden. So wäre eine Reihenfolge von Exon 5,1,4,2,3 möglich. Es wäre aber auch denkbar, dass nicht alle Exons Verwendung finden würden, z. B. würde für ein Genprodukt nur die Exone 1,2 und 3 benötigt. Dieses alternative Spleißen ermöglicht, dass man zwar nur 20.000 Gene hat, aber dennoch wesentlich mehr Proteine bilden kann. Hierbei handelt es sich aber um einen zielgerichteten Vorgang, das sowohl durch regulatorische Elemente in der RNA-Sequenz, an die spezifische Proteine binden, als auch durch Sekundärstrukturen reguliert wird, die sich mittels Hybridisierung von Nucleotiden im einzelsträngigen RNA-Molekül bilden. Wie aktuelle Untersuchungen zeigen, unterliegt etwa die Hälfte aller menschlichen Gene einem alternativen Spleißen. Mithilfe des alternativen Spleißens lassen sich möglicherweise auch unterschiedliche Komplexitätsniveaus von Lebewesen erklären.

Obwohl beispielsweise die Genome von Mensch und Schimpanse etwa gleich groß sind, kommt es im menschlichen Gehirn zu mehr alternativen Spleißvorgängen als beim Schimpansen.

Wir haben verschiedene Typen der RNA kennengelernt: mRNA, tRNA und rRNA. Es gibt jedoch eine Reihe weiterer RNA-Moleküle, die teilweise von den nicht-proteincodierenden DNA-Sequenzen exprimiert werden. Die RNA-Produkte dieser Regionen sind häufig sehr klein und deshalb schwer nachzuweisen. Man bezeichnet die winzigen RNA-Moleküle bei Prokaryoten und Eukaryoten als mikroRNA (miRNA, Abb. 116).

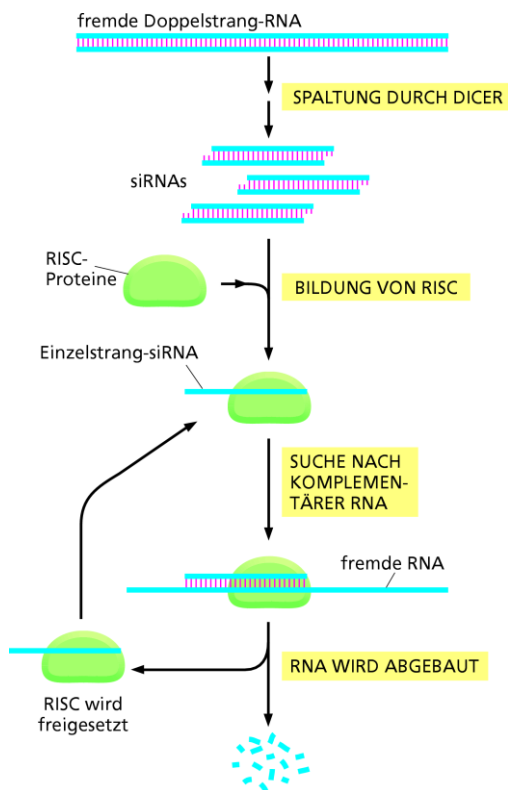


© 2012 Wiley-VCH, Weinheim
 Alberts – Lehrbuch der Molekularen Zellbiologie
 ISBN: 978-3-527-32624-6 Fig. 08-026

Abbildung 116 miRNA. Die miRNA-Vorstufe wird weiter bearbeitet und es entsteht eine reife miRNA. Diese lagert sich mit einer Gruppe von Proteinen zu einem Komplex zusammen, den man RISC nennt. Die miRNA leitet den RISC dann zu mRNAs mit einer komplementären Nukleotidsequenz. Je nachdem, wie groß der komplementäre Bereich ist, wird die Ziel-mRNA entweder schnell von einer Nuklease des RISC abgebaut oder an einen Ort im Cytoplasma transportiert, an dem andere zelluläre Nukleasen sie zerstören.

Die internationale miRNA-Datenbank „miRBase“ hatte im Frühjahr 2018 rund 30.000 Einträge, davon bezogen sich etwa 2660 auf das menschliche Genom. Diese miRNAs sind etwa 22 Basen lang und es gibt für jede mehrere Ziel-mRNAs. miRNAs werden in Form längerer Vorstufen transkribiert, die sich zu doppelsträngigen mRNA-Molekülen falten und dann durch eine Reihe von Reaktionsschritten zu einzelsträngigen miRNAs prozessiert werden. Ein Proteinkomplex lenkt die miRNA zu ihrer Ziel-mRNA, wo die Translation blockiert wird. Diese bemerkenswerte Konservierung des miRNA-vermittelten Gen-Silencings deutet darauf hin, dass es sich um einen evolutionär sehr alten Mechanismus von großer biologischer Bedeutung handelt.

Neben den miRNAs gibt es eine weitere Art von RNA-Molekülen, die auf ähnliche Weise wirken: kleine interferierende RNAs (small interfering RNAs, siRNAs, Abb. 117). Diese treten häufig bei Virusinfektionen auf, wenn zwei komplementäre Stränge eines Virusgenoms transkribiert werden. Dabei entstehen große doppelsträngige RNAs, die dann wie im Fall der miRNAs in kürzere, einzelsträngige Sequenzen umgewandelt werden.



© 2012 Wiley-VCH, Weinheim
 Alberts - Lehrbuch der Molekularen Zellbiologie
 ISBN: 978-3-527-32824-6 Fig. 08-027

Abbildung 117 siRNA. siRNAs zerstören fremde RNAs. Doppelsträngige RNAs von einem Virus oder einem transponierbaren genetischen Element werden zuerst von einer Nuklease gespalten, die Dicer genannt wird. Die so entstandenen doppelsträngigen Fragmente werden in RISCs aufgenommen. Im RISC wird ein Strang der Doppelhelix verworfen und der andere Strang wird verwendet, um komplementäre RNAs aufzuspüren und zu zerstören. Dieser Mechanismus bildet die Grundlage für die RNA-Interferenz (RNAi).

Diese binden an die Ziel-RNA und verursachen deren Abbau. Kleine interferierende RNAs leiten sich auch aus Transposonsequenzen ab, die in den Genomen der Eukaryoten weit verbreitet sind. Demnach hat sich wahrscheinlich das Abschalten von Genen durch siRNAs in der Evolution als Abwehrmechanismus gegen die Translation von Virus- und Transposonsequenzen entwickelt. miRNAs und siRNAs sind ähnliche Moleküle, die von den gleichen zellulären Enzymen prozessiert werden. Es gibt jedoch einen wichtigen Unterschied:

miRNAs werden von DNA-Sequenzen synthetisiert, die von den Zielsequenzen getrennt sind, während siRNAs gegen die Sequenz gerichtet sind, aus der sie stammen.

Die Bildung einer mRNA führt übrigens nicht automatisch zur Translation. Oft wird wesentlich mehr RNA gebildet als Proteine tatsächlich vorhanden sind. Auch der umgekehrte Fall kann eintreten: es gibt mehr Proteine als mRNA. Es gibt also Prozesse in der Zelle, die die Translation regulieren, sowie die Verweildauer eines Proteins in der Zelle.

Es gibt eine Reihe verschiedener Mechanismen, durch welche die Translation der mRNA reguliert werden kann. Zum einen kann die Translation durch siRNAs und miRNAs blockiert werden, zum anderen kann die Cap-Struktur der RNA beeinflusst werden. Dies ist ein chemisch modifiziertes Guanosintriphosphat (GTP). Dadurch wird

später bei der Translation die Bindung der mRNA an das Ribosom unterstützt und die mRNA ist vor einem Abbau geschützt. Eine mRNA, die ein nichtmodifiziertes GTP-Molekül als Cap-Struktur enthält, wird nicht translatiert. In einem anderen System blockieren Repressorproteine die Translation direkt. So bindet beispielsweise das Protein Ferritin in Säugerzellen freie Eisen(II)-Ionen (Fe^{2+}). Wenn Eisen im Überschuss vorhanden ist, nimmt die Ferritinsynthese erheblich zu, aber die Menge der Ferritin-mRNA bleibt konstant (Abb. 118).

Post-transcriptional regulation of ferritin and transferrin receptor

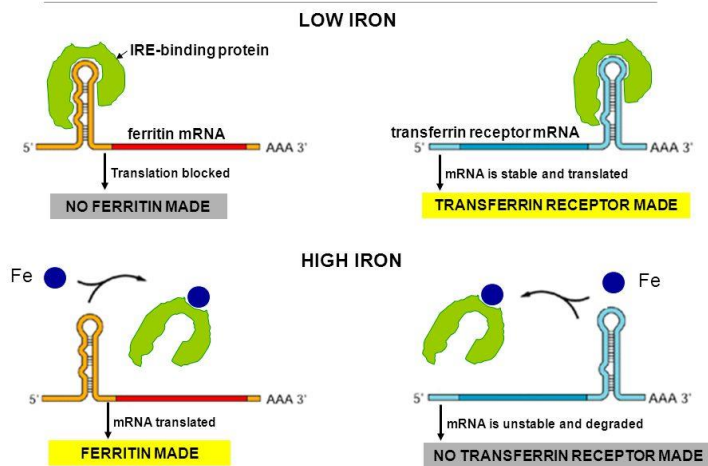


Abbildung 118 Regulation der Ferritin-Synthese

Offenbar ist die Zunahme der Ferritinsynthese auf eine erhöhte Translationsrate der mRNA zurückzuführen. Wenn der Eisenspiegel in einer Zelle gering ist, bindet ein Repressorprotein der Translation an die nichtcodierende 5'-Region der Ferritin-mRNA und hemmt die Translation, indem es die Bindung an das Ribosom verhindert. Wenn der Eisenspiegel steigt, bindet ein Teil der überschüssigen Fe^{2+} -Ionen an den Repressor, dessen dreidimensionale Struktur sich daraufhin ändert, sodass er sich von der mRNA ablöst und die Translation voranschreiten kann.

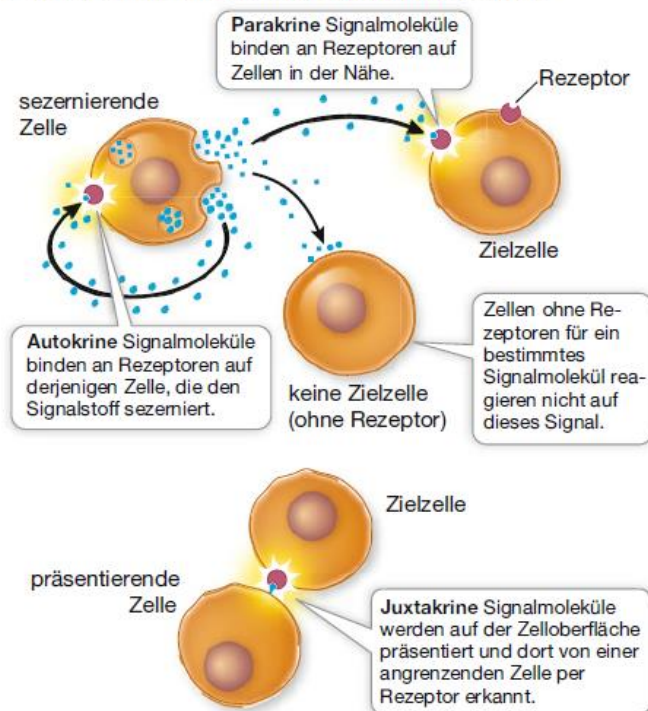
Kapitel 12: Zellkommunikation

Alle Zellen verarbeiten ihre Informationen aus ihrer Umwelt. Ein Signal kann aus der Umgebung des Organismus stammen, wie etwa der Geruch einer weiblichen Motte, die bei Dunkelheit einen Paarungspartner anlockt. Ein Signal kann aber auch von innerhalb des Organismus kommen, wie bei Zellen der Leber, wo von Zellen anderer Organe gebildete Signalmoleküle die Aufnahme oder die Freisetzung von Glucose steuern. Natürlich bedeutet das reine Vorhandensein eines Signals noch nicht, dass eine Zelle darauf reagiert. Damit eine Zelle auf ein Signal reagieren kann, benötigt sie ein spezifisches Rezeptorprotein, welches das Signal aufnimmt; außerdem ist ein Signalübertragungsweg (Signaltransduktionsweg) nötig, damit die aufgenommene Information zelluläre Prozesse auch beeinflussen kann. Ein Signaltransduktionsweg ist eine Abfolge von molekularen Vorgängen und chemischen Reaktionen, die zu einer Reaktion bzw. einer Antwort der Zelle auf ein Signal führen. Signaltransduktionswege variieren im Detail sehr stark, ihre Grundstruktur ist ihnen jedoch gemeinsam: Ein solcher Weg beginnt mit einem Signal und dessen spezifischem Rezeptor und mündet in eine Reaktion. Signalmoleküle lassen sich in drei Typen einteilen (Abb. 119):

1. Autokrine Signalmoleküle wirken auf diejenigen Zellen, die sie herstellen. Beispielsweise wachsen viele Tumoren so schnell, weil sie die entsprechenden Signalmoleküle selbst produzieren und so ihre eigenen Zellteilungen stimulieren.
2. Parakrine Signalmoleküle diffundieren zu Zellen in der Nähe, die sie beeinflussen. Ein Beispiel hierfür tritt im Zuge der Entzündungsreaktion bei einer Verletzung der Haut auf. Signalmoleküle aus den Hautzellen erreichen bestimmte Blutzellen in der Nähe, die dann den Heilungsprozess unterstützen.
3. Endokrine Signalmoleküle gelangen bei Tieren mit dem Blutkreislauf zu weit entfernten Zellen, die sie beeinflussen. Hierzu zählen z. B. Hormone.

Daneben gibt es noch juxtakrine Signalmoleküle. Sie bleiben in der Plasmamembran der produzierenden Zelle verankert und beeinflussen ausschließlich Zellen, die mit der signalproduzierenden Zelle in direktem Kontakt stehen. Diese Art der Signalübertragung ist insbesondere während der Entwicklung häufig anzutreffen, wenn die Zellen Gruppen bilden und sich spezialisieren.

a Wirkung von Signalmolekülen in räumlicher Nähe



b Wirkung von Signalmolekülen in weiterer Entfernung

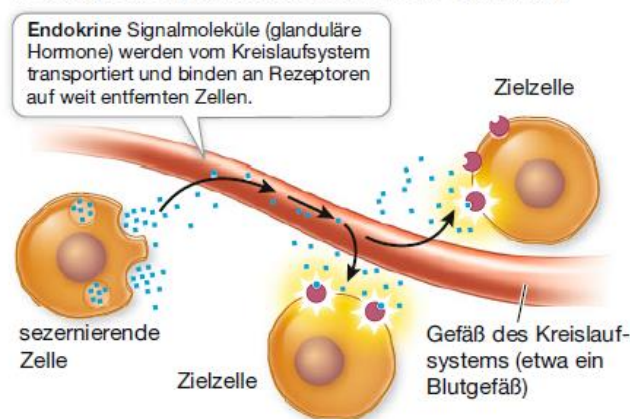


Abbildung 119 Formen der Zellkommunikation

Liganden

Damit die Information durch ein Signal auf eine Zielzelle übertragen werden kann, muss diese das Signal empfangen und darauf reagieren können. Das ist die Aufgabe des Rezeptors. In einem vielzelligen Lebewesen können alle Zellen sämtlichen chemischen Signalen ausgesetzt sein, aber die meisten von ihnen sind nicht in der Lage, auf diese Signale zu reagieren. Nur Zellen mit dem erforderlichen Rezeptor können eine Reaktion zeigen. Diese Spezifität der Bindung stellt sicher, dass nur diejenigen Zellen, die einen spezifischen Rezeptor synthetisieren (exprimieren), auf ein bestimmtes Signal reagieren.

An dieser Reaktion kann ein Enzym beteiligt sein, das eine biochemische Reaktion katalysiert, oder auch Transkriptionsfaktoren, die die Expression von bestimmten Genen an- und abschalten können.

Rezeptoren, die ein Signalmolekül erkennen, haben eine spezifische Bindungsstelle für dieses Molekül. Diese Signalmoleküle werden auch Ligand genannt (Abb. 120). Die Bindung eines Liganden führt zu einer Veränderung der dreidimensionalen Struktur des Rezeptorproteins, und diese Konformationsänderung löst eine zelluläre Reaktion aus. Der Ligand trägt zu dieser Reaktion nicht weiter bei. Tatsächlich wird der Ligand normalerweise nicht verändert. Seine Funktion besteht ausschließlich darin, „an die Tür zu klopfen“. Die Empfindlichkeit einer Zelle für ein Signalmolekül wird zum Teil von der Affinität des Rezeptors für den Liganden bestimmt, also der Wahrscheinlichkeit, mit der der Rezeptor den Liganden bei einer gegebenen Konzentration bindet. Diese Bindung ist reversibel, d. h. das Signalmolekül bleibt nicht dauerhaft an den Liganden gebunden. Würde das passieren, dann würde der Rezeptor ständig stimuliert und zelluläre Reaktionen auslösen.

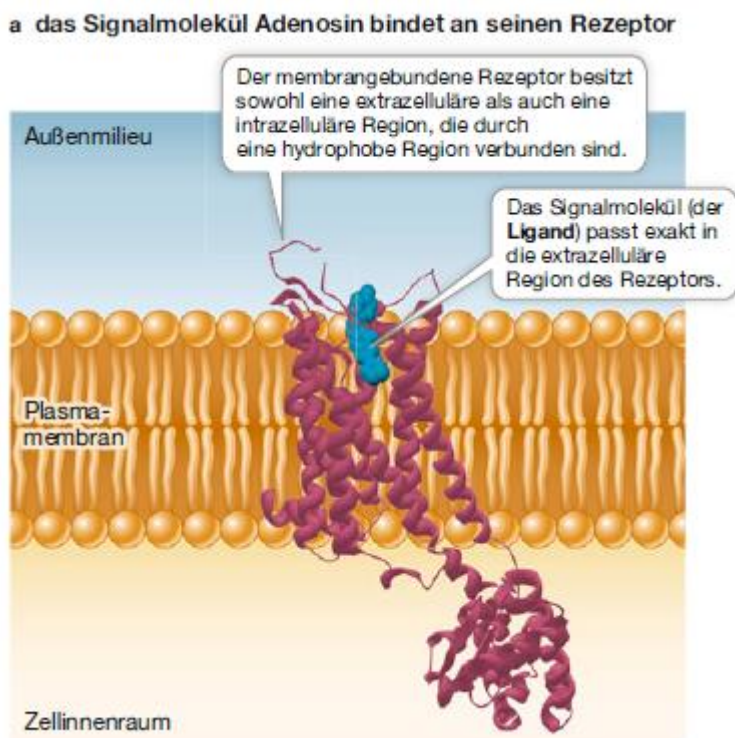


Abbildung 120 Rezeptor und Ligand

Bindet der Ligand an das Rezeptormolekül, kommt es bei letzterem zu einer Konformationsänderung. Durch die Konformationsänderung wird zum Beispiel eine zuvor im Inneren des Rezeptors verborgene Gruppe von Aminosäuren so exponiert, dass deren Seitenketten an einer biochemischen Reaktion teilnehmen können. Diese Reaktion kann die Bindung eines weiteren Moleküls sein, zum Beispiel eines anderen Proteins oder auch eines Enzymsubstrats.

Anstelle des normalen Liganden können auch andere Substanzen, die dem Liganden stark ähneln, an den Rezeptor binden. Agonisten sind chemische Verbindungen, die wie der Ligand die Signalübertragung durch den Rezeptor auslösen. Im Gegensatz

dazu binden Antagonisten, auch als Inhibitoren bezeichnet, an den Rezeptor, „fixieren“ seine Raumstruktur und hindern den natürlichen Liganden an einer Bindung; sie lösen jedoch keine Signalübertragung aus. Agonisten und Antagonisten können natürlicher Herkunft sein oder sie werden gezielt im Labor entwickelt.

Viele Substanzen, die das menschliche Verhalten beeinflussen, binden an spezifische Rezeptoren im Gehirn und verhindern so die Bindung des eigentlichen Liganden des jeweiligen Rezeptors. Ein Beispiel dafür ist das Koffein, das wahrscheinlich weltweit am häufigsten konsumierte Aufputschmittel. Im Gehirn fungiert das Nucleosid Adenosin als natürlicher Ligand, der an einen Rezeptor auf Nervenzellen bindet. Dadurch wird ein Signaltransduktionsweg ausgelöst, der die Gehirnaktivität reduziert, insbesondere das Gefühl aktiver Wachheit. Da Koffein eine ähnliche Molekülstruktur besitzt wie Adenosin, bindet es ebenfalls an den Adenosinrezeptor (Abb. 121). Im Falle einer Bindung wird der Signaltransduktionsweg jedoch nicht aktiviert. Stattdessen wird der Rezeptor blockiert und die Bindung von Adenosin verhindert. Die Nervenzellen bleiben weiterhin aktiv und der Wachzustand bleibt fühlbar bestehen.

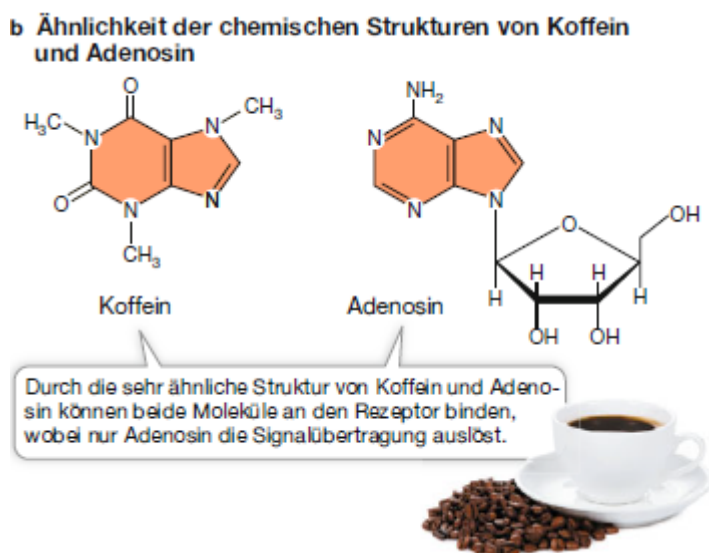


Abbildung 121 Koffein als Ligand

Rezeptoren

Die im Körper aktiven Signalmoleküle sind sehr vielfältig. Einige Liganden sind hydrophob (unpolar) und können durch Membranen diffundieren, andere können dieses nicht. Physikalische Reize wie Licht unterscheiden sich in ihrer Fähigkeit, in Zellen oder Gewebe einzudringen. Entsprechend kann man Rezeptoren nach ihrer Lokalisation in der Zelle klassifizieren (Abb. 122); diese hängt stark von der Art seines Liganden bzw. Reizes ab:

1. Membranrezeptoren: Große oder polare Liganden können die Lipiddoppelschicht nicht passieren. Insulin zum Beispiel ist ein Proteohormon, das nicht durch die Plasmamembran diffundieren kann. Stattdessen bindet das

Molekül an einen Transmembranrezeptor mit einer extrazellulären Bindungsdomäne.

2. intrazelluläre Rezeptoren: Kleine oder unpolare Liganden können durch die unpolare Phospholipiddoppelschicht der Plasmamembran diffundieren und so in die Zelle gelangen. Das Steroidhormon Östrogen beispielsweise ist lipidlöslich und kann daher durch die Plasmamembran diffundieren. Es bindet an einen intrazellulären Rezeptor. Licht bestimmter Wellenlängen kann relativ ungehindert in die Zellen eines Blattes vordringen, sodass zahlreiche Lichtrezeptoren von Pflanzen ebenfalls intrazellulär lokalisiert sind.

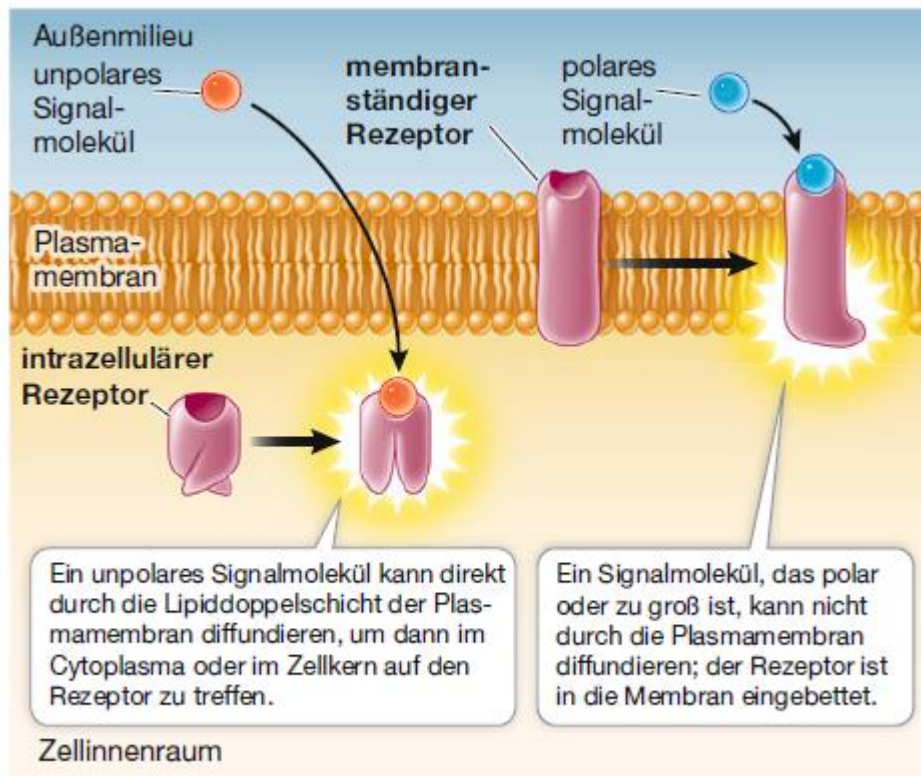


Abbildung 122 Membranrezeptoren und intrazelluläre Rezeptoren

Bei komplexen Eukaryoten wie den Wirbeltieren oder Landpflanzen kennt man drei genau untersuchte Gruppen von Plasmamembranrezeptoren, die man entsprechend ihrer Funktionen unterteilt hat:

Ligandengesteuerte und andere Ionenkanäle (Abb. 123): Wie wir erfahren haben, enthält die Plasmamembran vieler Zelltypen ligandengesteuerte Ionenkanäle, die Ionen wie Na^+ , K^+ , Ca^{2+} oder Cl^- in die Zelle hinein oder aus ihr heraus diffundieren lassen. Der Öffnungsmechanismus besteht in der Veränderung der dreidimensionalen Struktur des Kanalproteins nach der Wechselwirkung mit dem Liganden. Ein ligandengesteuerter Ionenkanal reagiert also auf einen spezifischen chemischen Signalstoff. Bei anderen Ionenkanälen kann dieses Signal ein sensorischer Reiz wie Licht oder Schall sein, und bei spannungsgesteuerten Ionenkanälen ist es die Veränderung der elektrischen Spannung über der Membran. Der Acetylcholinrezeptor, der bei Wirbeltieren in der Plasmamembran der Skelettmuskelzellen vorkommt, ist ein

Beispiel für einen ligandengesteuerten Ionenkanal. Dabei handelt es sich um einen Natriumkanal, der den Liganden Acetylcholin bindet. Acetylcholin ist ein sogenannter Neurotransmitter – ein Signalmolekül, das von Nervenzellen freigesetzt wird. Wenn zwei Moleküle Acetylcholin an den Rezeptor binden, öffnet er den Kanal für etwa eine Tausendstelsekunde. Diese Zeit genügt, damit Na^+ , das außerhalb der Zelle viel höher konzentriert ist als innen, in die Zelle strömen kann. Die treibende Kraft ist sowohl der Konzentrationsgradient als auch die elektrische Potenzialdifferenz (elektrischer Gradient). Die Veränderung der Na^+ -Konzentration führt über ein Muskelaktionspotenzial schließlich zur Kontraktion des Muskels.

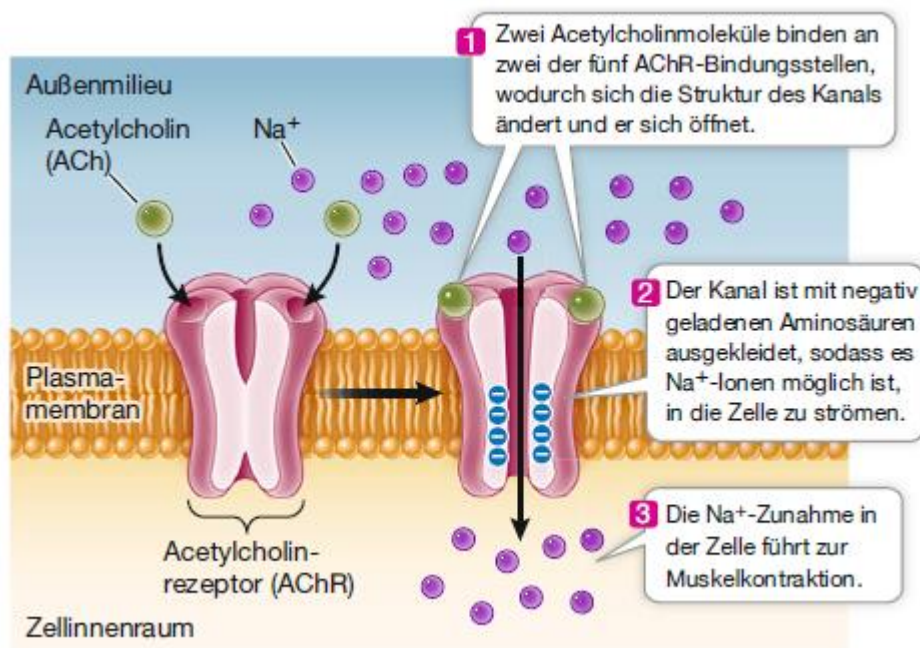


Abbildung 123 Ionenkanäle

Enzymgekoppelte Rezeptoren, insbesondere Proteinkinaserezeptoren (Abb. 124). Bestimmte eukaryotische Rezeptorproteine sind mit einem Enzym gekoppelt. Die wichtigste Gruppe solcher Enzyme sind die Proteinkinasen; die damit gekoppelten Rezeptoren heißen entsprechend Proteinkinase-Rezeptoren. Proteinkinasen katalysieren die eigene Phosphorylierung (Addition von Phosphat) oder die Phosphorylierung anderer Proteine, was die Proteinstruktur und damit die Proteinfunktion verändert. Die Phosphorylierung von Proteinen hat in der Biologie eine herausragende Bedeutung. Beim Menschen codieren z. B. 500 Gene für Proteinkinasen. Proteinkinasen wirken also an unzähligen Stellen in unserem Körper und beeinflussen damit auf die eine oder andere Weise unsere sämtlichen Lebensäußerungen. Der Insulinrezeptor ist ein Beispiel für einen Proteinkinaserezeptor. Insulin ist ein Proteohormon, das in der Bauchspeicheldrüse produziert wird. Der zugehörige Rezeptor besteht aus jeweils zwei Kopien zweier verschiedener Proteinuntereinheiten, alpha und beta. Wenn Insulin an den Rezeptor bindet, wird dieser aktiviert und kann sich selbst und bestimmte Proteine im Cytoplasma phosphorylieren, die man treffend als Substrate der Insulinreaktion bezeichnet. Diese Proteine setzen dann zahlreiche zelluläre Reaktionen in Gang, beispielsweise das Einfügen von Glucosetransportern in die Plasmamembran.

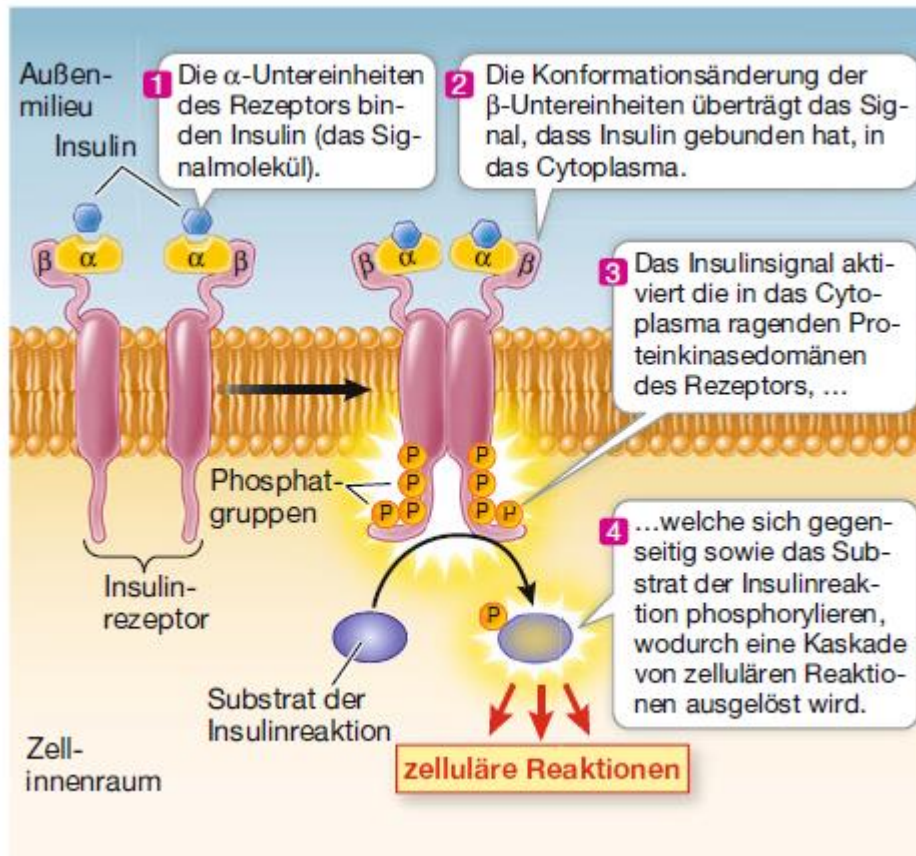


Abbildung 124 Enzymgekoppelte Rezeptoren

G-Protein-gekoppelte Rezeptoren (Abb. 125). Diese Rezeptoren sind an der Verarbeitung eines breiten Spektrums an Signalen beteiligt, darunter Licht, das auf die Netzhaut von Wirbeltieren trifft (Photorezeptoren), und Duftstoffe (Geruchsrezeptoren). Außerdem spielen sie bei der Regulation von Stimmungen und Verhalten (wie das Paarungsverhalten von Säugetieren) eine Rolle. So zählen die Rezeptoren der Hormone Oxytocin und ADH, die das Paarungsverhalten der Wühlmäuse steuern zu den G-Protein-gekoppelten Rezeptoren. Die sieben Transmembranhelices dieser Rezeptoren durchqueren die Phospholipiddoppelschicht und sind über kurze Schleifen miteinander verbunden, die sich dadurch abwechselnd außerhalb oder innerhalb der Zelle befinden. Die Bindung eines Liganden an die extrazelluläre Region des Rezeptors verändert die Struktur der cytoplasmatischen Region, sodass dort eine Stelle zugänglich wird, an die ein mobiles Membranprotein, ein G-Protein, binden kann. Das G-Protein taucht teilweise in die Lipiddoppelschicht ein. Viele G-Proteine bestehen aus drei Proteinuntereinheiten und können drei verschiedene Moleküle binden: den Rezeptor, ein GDP bzw. GTP (Guanosindiphosphat bzw. -triphosphat) und an ein Effektorprotein. Wenn das G-Protein an ein aktiviertes Rezeptorprotein bindet, tritt GTP an die Stelle von GDP. Normalerweise wird gleichzeitig an der extrazellulären Seite des Rezeptors der Ligand freigesetzt. Die Bindung von GTP verursacht eine Konformationsänderung des G-Proteins. Die Untereinheit mit dem gebundenen GTP trennt sich von dem übrigen G-Protein und diffundiert in der Ebene der Phospholipiddoppelschicht, bis sie auf ein Effektorprotein trifft, an das sie binden kann. Die Bindung der G-Protein-Untereinheit mit dem gekoppelten GTP aktiviert den Effektor – der zum Beispiel ein Enzym oder ein

Ionenkanal sein kann – und führt so zu Veränderungen der Zellfunktion. Nach der Bindung und Aktivierung des Effektorproteins löst sich die G-Protein-Untereinheit und kann weitere Effektorproteinmoleküle aktivieren, bis die intrinsische GTPase-Aktivität der Untereinheit das GTP zu GDP hydrolysiert und so inaktiviert. Die G-Protein-Untereinheit kann nun wieder an die beiden anderen G-Protein-Untereinheiten binden. Sobald sich die drei Untereinheiten zusammengefügt haben, kann das G-Protein erneut an einen aktivierten Rezeptor binden. Nach der Bindung tauscht der aktivierte Rezeptor das GDP auf dem G-Protein gegen GTP aus und der ganze Zyklus beginnt von vorne.

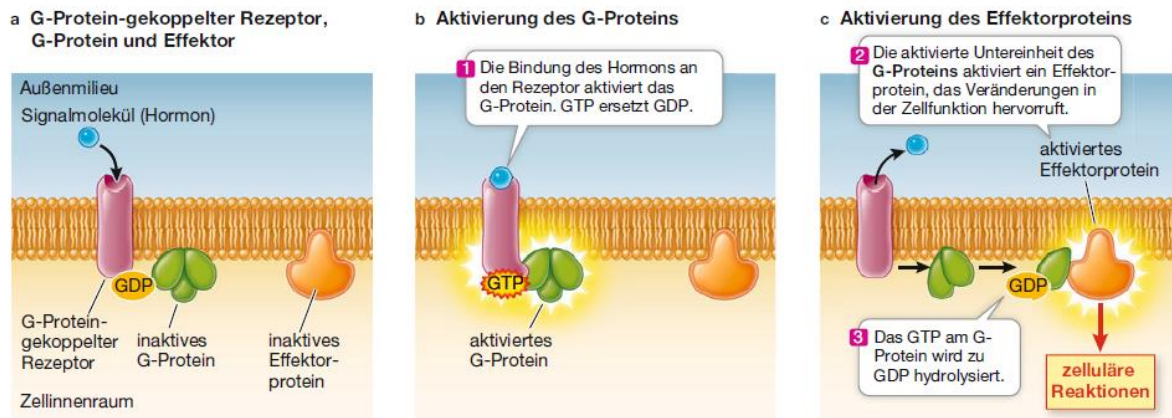


Abbildung 125 G-Protein gekoppelter Rezeptor

Intrazelluläre Rezeptoren (Abb. 126) sind in der Zelle lokalisiert und dienen dem Erkennen von physikalischen Reizen wie Licht (z. B. einige Photorezeptoren bei Pflanzen) oder von Signalmolekülen, die durch die Plasmamembran diffundieren können (z. B. Steroidhormone bei Tieren). Bei zahlreichen intrazellulären Rezeptoren handelt es sich um Transkriptionsfaktoren, von denen wiederum einige bis zu ihrer Aktivierung im Cytoplasma lokalisiert sind. Nach Bindung ihres Liganden diffundieren diese Transkriptionsfaktoren durch die Kernporen in den Zellkern, wo sie an die DNA binden und die Expression von bestimmten Genen modulieren. Ein typisches Beispiel dafür ist der Rezeptor des Steroidhormons Cortisol. Der Rezeptor ist normalerweise an ein Chaperonprotein gebunden, das die Diffusion des Rezeptors in den Zellkern verhindert. Durch die Bindung eines Hormons verändert der Rezeptor seine Struktur, sodass er sich vom Chaperon löst. Dies ermöglicht es dem Rezeptor, in den Zellkern zu diffundieren, wo er die Transkription der DNA beeinflusst. Eine andere Gruppe von intrazellulären Rezeptoren ist stets im Zellkern lokalisiert, sodass ihre Liganden bis dorthin gelangen müssen, bevor sie an den Rezeptor binden können.

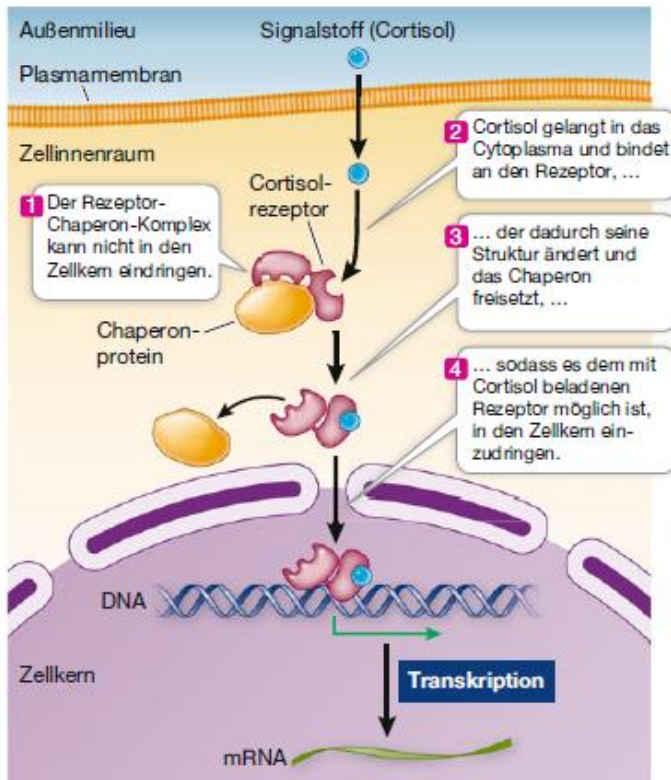
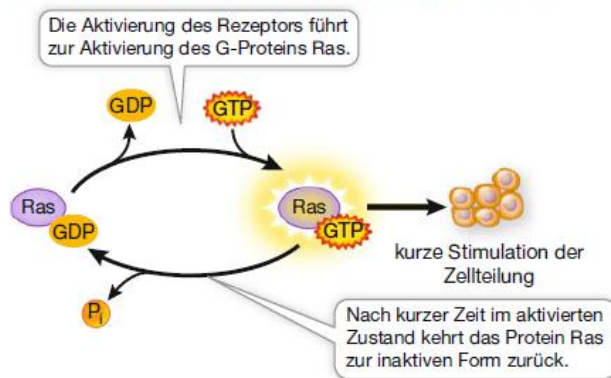


Abbildung 126 Intrazellulärer Rezeptor

Signalkaskaden und sekundäre Botenstoffe

Es gibt verschiedene Arten von Signalen und Rezeptoren und wie zu erwarten ist, variieren auch die Wege, auf denen die Signale übertragen werden, wie auch die ausgelösten zellulären Reaktionen. Einige Signaltransduktionswege sind verhältnismäßig einfach und direkt, andere hingegen umfassen zahlreiche Schritte. Daran können Enzyme und auch Transkriptionsfaktoren beteiligt sein. Außerdem können weitere Moleküle, die man als sekundäre Botenstoffe (second messenger) nennt, durch das Cytoplasma diffundieren und Zwischenschritte im Signalübertragungsweg vermitteln. In vielen Fällen kann ein Signal eine Kaskade von Ereignissen auslösen, bei denen Proteine mit anderen Proteinen in Wechselwirkung treten, die wiederum mit weiteren Proteinen interagieren, bis die letztendlichen Reaktionen ablaufen. Über solch eine auf Proteinen basierende Signaltransduktionskaskade kann ein ursprüngliches Signal in der Zelle sowohl verstärkt als auch verteilt werden, um so in der Zielzelle verschiedene Reaktionen hervorzurufen. Wir haben uns z. B. die G-Protein-gekoppelten Rezeptoren angeschaut, bei dem ein G-Protein solch eine Signaltransduktionskaskade auslöst (Abb. 6). G-Proteine fungieren als eine Art Schalter, die sich im „An“- oder auch im „Aus“-Zustand befinden können. Stehen sie auf „Aus“, dann ist GDP an das G-Protein gebunden, im „An“-Zustand ist das Protein mit GTP beladen. Eines der bekannten G-Proteine heißt Ras, welches eine Reihe von Ereignissen in Gang setzt, bei denen eine Proteinkinase die nächste aktiviert und so weiter (Abb. 127).

a die Funktion des Proteins Ras in einer normalen Zelle



b die Funktion von Ras in einer Tumorzelle

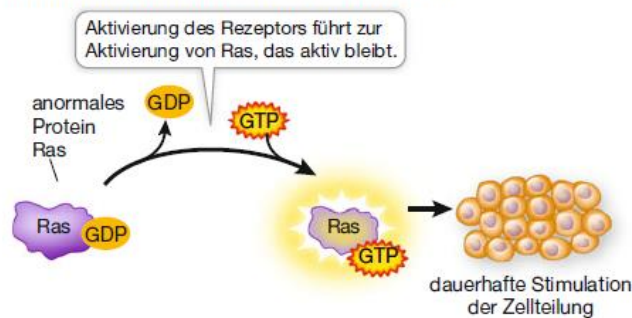


Abbildung 127 Das G-Protein Ras

Einen solchen Vorgang bezeichnet man als Proteinkinasekaskade (Abb. 128). Derartige Kaskaden sind zentraler Bestandteil der Regulation zahlreicher zellulärer Aktivitäten.

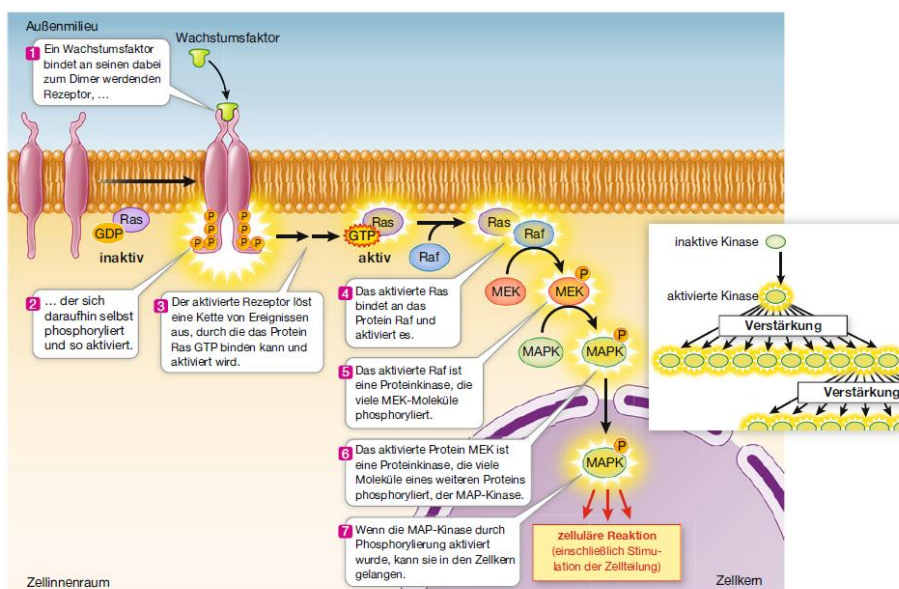


Abbildung 128 Eine Proteinkinasekaskade. Bei einer Proteinkinasekaskade wird eine Reihe von Proteinen nacheinander aktiviert.

Proteinkinasekaskaden sind aus vier Gründen nützliche Signalüberträger:

1. Bei jedem Schritt der Kaskade erfolgt eine Verstärkung des Signals, da jedes neu aktivierte Proteinkinase-Molekül ein Enzym ist, das die Phosphorylierung von zahlreichen Zielproteinmolekülen katalysieren kann.
2. Die Information, dass ein bestimmtes Signal an der Plasmamembran angekommen ist, wird in den Zellkern übertragen, wo häufig die Expression zahlreicher Gene moduliert wird.
3. Die Vielzahl der Schritte gewährleistet eine Spezifität des Vorgangs.
4. Unterschiedliche Zielproteine in den einzelnen Schritten der Kaskade erlauben eine Variation der Reaktion.

In vielen Fällen existiert ein kleines Molekül, das zwischen den aktivierten Rezeptor und den Beginn der eigentlichen Kaskade geschaltet ist. Ein solches Molekül wurde entdeckt, als man die Aktivierung des Leberenzym Glykogen-Phosphorylase durch das Hormon Adrenalin untersuchte. Dieses Hormon wird freigesetzt, wenn sich ein Tier in einer Stresssituation befindet und schnell Energie für die Kampf-oder-Flucht-Reaktion verfügbar sein muss. Die Glykogen-Phosphorylase katalysiert den Abbau von Glykogen, das in Leberzellen gespeichert ist, sodass die entstehenden Glucosemoleküle in das Blut freigesetzt werden können. Das Enzym kommt im Cytoplasma der Leberzellen vor, ist jedoch inaktiv, wenn kein Adrenalin vorhanden ist.

Wie sich herausstellte, kann das Adrenalin die Glykogen-Phosphorylase auch in Leberzellextrakten aktivieren, aber nur dann, wenn der gesamte Zellinhalt inklusive Fragmenten der Plasmamembran vorhanden ist. Unter diesen Bedingungen war Adrenalin an die Fragmente der Plasmamembranen gebunden (wo der Adrenalinrezeptor lokalisiert ist), aber die aktivierte Phosphorylase befand sich frei in der Lösung. Die Zugabe von Adrenalin zum Cytoplasma mit der inaktiven Phosphorylase führte nicht zu einer Aktivierung des Enzyms. Daraus konnte abgeleitet werden, dass das Signal von Adrenalin über einen sekundären Botenstoff (second messenger) übertragen wird. Dieser sekundäre Botenstoff ist das zyklische AMP (cAMP). Es wird von dem Enzym Adenylat-Cyclase aus ATP gebildet. Die Adenylat-Cyclase wird vom G-Protein-gekoppelten Adrenalinrezeptor aktiviert (Abb. 129).

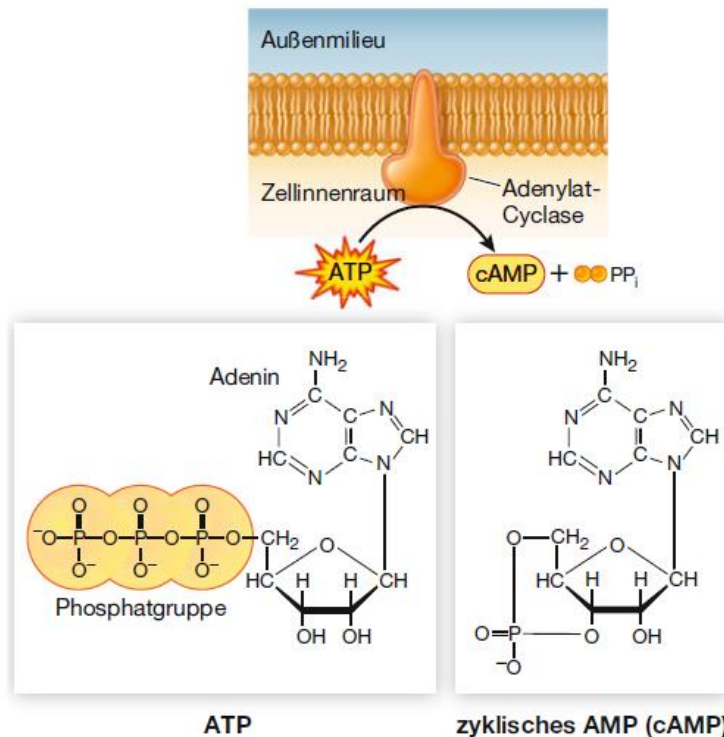


Abbildung 129 Die Bildung von zyklischem AMP. Die Bildung von cAMP aus ATP wird von der Adenylat-Cyclase katalysiert, die durch G-Proteine aktiviert wird

Im Gegensatz zur hohen Spezifität der Rezeptorbindung erlauben es sekundäre Botenstoffe wie cAMP der Zelle, auf ein einzelnes Ereignis an der Plasmamembran mit einer Vielzahl an zellulären Ereignissen innerhalb der Zelle zu reagieren. Sekundäre Botenstoffe dienen also dazu, das Signal rasch zu verstärken und zu verbreiten. So führt beispielsweise die Bindung eines einzigen Adrenalinmoleküls zur Produktion von zahlreichen Molekülen cAMP, die dann ebenso viele Kopien des Zielenzyms aktivieren, indem sie nichtkovalent an diese binden.

Im Fall von Adrenalin und der Leberzelle ist die Glykogen-Phosphorylase zudem nur eines von mehreren Enzymen, die aktiviert werden.

Sekundäre Botenstoffe sind häufig an der Vernetzung von Signalwegen, dem Crosstalk, beteiligt. Die Aktivierung des Adrenalinrezeptors ist nicht der einzige für eine Zelle gangbare Weg, um cAMP herzustellen; es gibt zahlreiche Ziele für cAMP in der Zelle, die gleichzeitig Komponenten anderer Signalwege darstellen.

Auch einige Lipide sind als sekundäre Botenstoffe bekannt. Bei der Bildung bestimmter sekundärer Botenstoffe werden gewisse Phospholipide durch Enzyme, die man als Phospholipasen bezeichnet, durch hydrolytische Spaltung in ihre Bestandteile zerlegt.

Die am besten untersuchten Beispiele für von Lipiden abgeleitete sekundäre Botenstoffe hängen mit der Hydrolyse des Phospholipids Phosphat-Idylinositol-4,5-bisphosphat (PIP₂) zusammen. Wie alle Phospholipide hat PIP₂ einen hydrophoben Teil, der in die Plasmamembran eingebettet ist: zwei an Glycerol gebundene Fettsäureschwänze, zusammen als Diacylglycerol (DAG) bezeichnet. Der hydrophile Anteil von PIP₂, der in das Cytoplasma hineinragt, ist Inositoltrisphosphat (IP₃).

Wie bei cAMP sind die Rezeptoren, die zu diesem System der sekundären Botenstoffe gehören, häufig mit G-Proteinen gekoppelt. Eine G-Protein-Untereinheit wird durch den Rezeptor aktiviert, diffundiert dann innerhalb der Plasmamembran und aktiviert die Phospholipase C, ein Enzym, das sich in der Membran befindet. Dieses Enzym spaltet IP3 von PIP2 ab, sodass in der Lipiddoppelschicht Diacylglycerol (DAG) entsteht: Sowohl IP3 als auch DAG sind sekundäre Botenstoffe. Sie gehören zu unterschiedlichen Signalübertragungswegen, die einander unterstützen und beide die Proteinkinase C (PKC) aktivieren.

Auch Calciumionen können als sekundäre Botenstoffe dienen. Sie kommen in den meisten Zellen nur in geringen Mengen vor. Außerhalb der Zellen und innerhalb des endoplasmatischen Reticulums ist die Ca^{2+} -Konzentration oftmals viel höher. Ca^{2+} -Pumpen in der Plasmamembran und der ER-Membran halten die Konzentrationsunterschiede unter ATP-Verbrauch aufrecht, indem sie Ca^{2+} aus dem Cytosol befördern. Anders als cAMP und die von Lipiden abgeleiteten sekundären Botenstoffe kann die Zelle Ca^{2+} nicht herstellen, sondern muss diese Ionen aus der Umgebung aufnehmen. Der Ca^{2+} -Spiegel im Cytoplasma und in Organellen wird durch das Öffnen und Schließen von Ca^{2+} -Ionenkanälen und die Aktivität der membranständigen Ca^{2+} -Pumpen reguliert.

Zahlreiche Signalstoffe können Calciumkanäle öffnen, beispielsweise IP3. Das Eindringen des Spermienzellkerns in eine Eizelle ist ein sehr wichtiges Signal, das zum Öffnen zahlreicher Calciumkanäle führt. Dadurch kommt es zu beträchtlichen Veränderungen, die das nun befruchtete Ei auf die anstehenden Zellteilungen vorbereiten. Unabhängig davon, welches Signal zum Öffnen der Calciumkanäle geführt hat, das Ergebnis ist eine deutliche Zunahme der Ca^{2+} -Konzentration im Zellplasma, die innerhalb eines Sekundenbruchteils um das Hundertfache ansteigen kann. Diese Zunahme aktiviert die Proteinkinase C.

Kapitel 13: Mitose und Zellzyklus

Wir wissen nun wie unsere DNA vermehrt wird über einen Prozess, der sich Replikation nennt. Hier werden wir uns mit der Vermehrung von Zellen befassen.

Grundsätzlich: Alle Zellen stammen von Zellen ab. Die Zellteilung ist bei allen Lebewesen Grundvoraussetzung für die Reproduktion (Fortpflanzung) und bei vielzelligen Organismen zudem für Wachstum, Regeneration und Reparatur. Für eine jede Zellteilung sind vier Ereignisse erforderlich:

1. Das Reproduktionssignal. Ein intrazelluläres oder extrazelluläres Signal löst die Zellteilung aus.
2. Die Replikation der DNA. Das genetische Material der Zelle muss dupliziert werden, damit jede der beiden neuen Zellen über einen vollständigen und identischen Satz von Genen verfügt. Diesen Schritt haben wir bereits kennengelernt und werden ihn hier nicht weiter vertiefen. Schaut dazu nochmal den Beitrag zur Replikation der DNA an.
3. Die Segregation der replizierten DNA. Die replizierte DNA muss auf die beiden neuen Zellen gleichmäßig verteilt werden.
4. Die Cytokinese. Die Enzyme und Organellen der neuen Zellen müssen nachproduziert und der Plasmamembran muss neues Material hinzugefügt werden (auch der Zellwand bei Lebewesen, die eine solche besitzen), damit sich die beiden Zellen voneinander trennen können.

Zellzyklus

Wie weiß eine Zelle, dass sie sich teilen muss?

Bei den Prokaryoten wirken häufig äußere Signale wie ein Wechsel von Umweltbedingungen oder die Konzentration von Nährstoffen als Auslöser für die Zellteilung. Das Bakterium *Escherichia coli* teilt sich bei ausreichender Versorgung mit Kohlenhydraten und Mineralstoffen alle 20 min.

Anders als Prokaryoten teilen sich eukaryotische Zellen nicht ständig, sobald die Bedingungen in ihrer Umgebung günstig sind. Bei vielzelligen Eukaryoten sind die meisten Zellen sehr spezialisiert und teilen sich tatsächlich nur selten. In einem eukaryotischen Organismus hängen die Signale für die Zellteilung nicht mit der Umgebung einer einzelnen Zelle zusammen, sondern mit den Bedürfnissen und Anforderungen des gesamten Lebewesens.

Die Zellteilung bei Eukaryoten wird durch den Zellzyklus reguliert. Der Zellzyklus lässt sich untergliedern in Mitose/Cytokinese und Interphase. Während der Interphase ist der Zellkern im Lichtmikroskop sichtbar und es finden die charakteristischen Zellfunktionen statt wie die DNA-Replikation. Diese Phase der Zelle beginnt, wenn die Cytokinese abgeschlossen ist, und endet mit dem Beginn der Mitose (M). Die Dauer des Zellzyklus ist bei den einzelnen Zelltypen sehr unterschiedlich. In der frühen Embryonalphase kann der Zellzyklus nur 30 min umfassen, während sich schnell teilende Zellen bei einem erwachsenen Menschen für einen vollständigen Zellzyklus

etwa 24 h benötigen. Allgemein gilt, dass sich die Zellen die meiste Zeit in der Interphase befinden.

Der Zellzyklus lässt sich einteilen in eine G1-Phase, S-Phase, G2-Phase und M-Phase.

Während der G1-Phase besteht jedes Chromosom aus einem einzigen, nichtreplizierten DNA-Molekül, das an Histone gebundenen ist. Für die unterschiedliche Dauer des Zellzyklus bei den einzelnen Zelltypen ist vor allem die Dauer der G1-Phase verantwortlich. Im Übergang von der G1-Phase zur S-Phase wird die Zelle darauf festgelegt, die DNA zu replizieren und anschließend die Mitose durchzuführen.

Während der S-Phase (Synthesephase) kommt es zur Replikation der DNA. Jedes Chromosom wird verdoppelt, es besteht also aus zwei Schwesterchromatiden, die bis zur Mitose miteinander verbunden bleiben.

Während der G2-Phase bereitet sich die Zelle auf die Mitose vor – etwa indem sie die Proteinkomponenten synthetisiert, die im Verlauf der Mitose gebraucht werden.

Manche in der G1-Phase befindliche Zellen treten in einen inaktiven Ruhezustand des Zellzyklus ein, die G0-Phase. Solche Zellen können oft unter bestimmten äußeren Bedingungen wieder in die G1-Phase und damit in den Zellzyklus wechseln, etwa aufgrund eines extrazellulären Signals.

Bei einem Zellzyklus von 24 h dauern diese untergeordneten Phasen normalerweise 11 h (G1-Phase), 8 h (S-Phase) und 4 h (G2-Phase), wobei die noch fehlende 1 h auf die Mitose/Cytokinese (M-Phase) entfällt (Abb. 130).

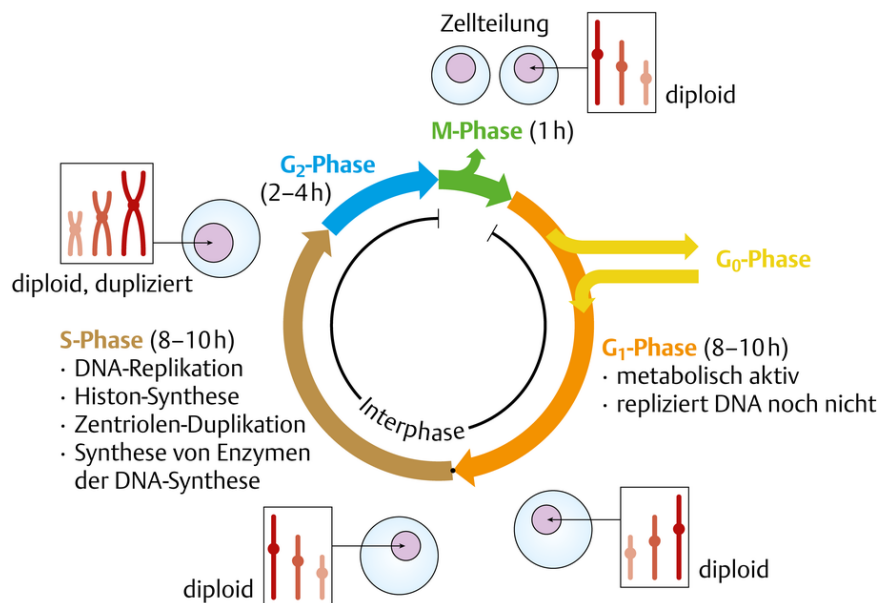


Abbildung 130 Zellzyklus

Spezifische Signale lösen die Ereignisse des Zellzyklus aus. Das Voranschreiten des Zellzyklus beruht auf den Aktivitäten der cyclinabhängigen Kinasen (Cdks). Das sind Enzyme, die die Übertragung einer Phosphatgruppe von ATP auf ein Zielprotein

katalysieren. Diese Phosphatgruppenübertragung bezeichnet man als Phosphorylierung. Das ändert die räumliche Struktur und damit die Aktivität des Proteins. Durch Katalyse der Phosphorylierung bestimmter Zielproteine spielen die Cdks an verschiedenen Stellen im Zellzyklus eine wichtige Rolle. Z. B.: Eine Arbeitsgruppe um James Maller untersuchte unreife Eier des Krallenfroschs *Xenopus laevis* und versuchte herauszufinden, wie sie zur Teilung und schließlich zur Bildung reifer Eier stimuliert werden. Man stellte fest, dass bei Einführen eines Proteins aus heranreifenden Eiern in unreife Eizellen diese zur Teilung angeregt wurden. Das Protein wurde als Reifungsfaktor (maturation promoting factor) bezeichnet. Gleichzeitig untersuchte Leland Hartwell mit seiner Gruppe den Zellzyklus der Bierhefe *Saccharomyces cerevisiae*. Dabei entdeckte er einen Stamm, der beim G1-S-Übergang angehalten wurde, da ihm eine Cdk fehlte.

Es stellte sich heraus, dass diese Cdk der Hefe und der Reifungsfaktor des Krallenfroschs ähnliche Eigenschaften besitzen. Weitere Untersuchungen bestätigten dann, dass das Protein der Froscheier tatsächlich eine Cdk ist. Man fand bald darauf bei vielen anderen Organismen, so auch beim Menschen, ähnliche Cdks, die den G1-S-Übergang regulieren.

Diese Kontrollstelle bezeichnet man heute als Restriktionspunkt (R-Punkt). Inzwischen sind weitere Cdks bekannt, die andere Abschnitte des Zellzyklus regulieren.

Cdks sind erst dann als Proteinkinasen enzymatisch aktiv, wenn sie an einen anderen Typ von Protein gebunden sind – an einen Aktivator mit der Bezeichnung Cyclin. Cyclin wurde von Tim Hunt bei Seeigeleiern entdeckt; danach konnte er zeigen, dass es auch bei Froscheiern wirkt. Die Bindung von Cyclin an die Cdk aktiviert die Cdk, indem es deren Konformation verändert, sodass das aktive Zentrum für die Substrate zugänglich wird. Der Cyclin/Cdk-Komplex, der den Übergang von der G1- zur S-Phase kontrolliert, ist nicht der einzige derartige Komplex, der bei der Regulation des eukaryotischen Zellzyklus eine Rolle spielt. Es gibt unterschiedliche Cyclin/Cdk-Komplexe, die aus unterschiedlichen Cyclinen und Cdks bestehen und in verschiedenen Zyklusphasen aktiv sind.

Als Beispiel dient der Cyclin/Cdk-Komplex, der den G1-S-Übergang kontrolliert. Der Cyclin/Cdk-Komplex katalysiert die Phosphorylierung des Retinoblastomproteins (RB). In vielen Zellen fungiert RB oder ein RB-ähnliches Protein als Inhibitor des Zellzyklus am R-Punkt. Um in die S-Phase einzutreten, muss eine Zelle die RB-Blockade umgehen. Hier liegt nun die Funktion des G1-S-Cyclin/Cdk-Komplexes. Der Komplex katalysiert das Anhängen einer Phosphatgruppe an RB. Dadurch verändert sich die dreidimensionale Struktur von RB, und es wird infolgedessen inaktiviert. Sobald die Wirkung von RB aufgehoben ist, kann der Zellzyklus voranschreiten (Abb. 131).

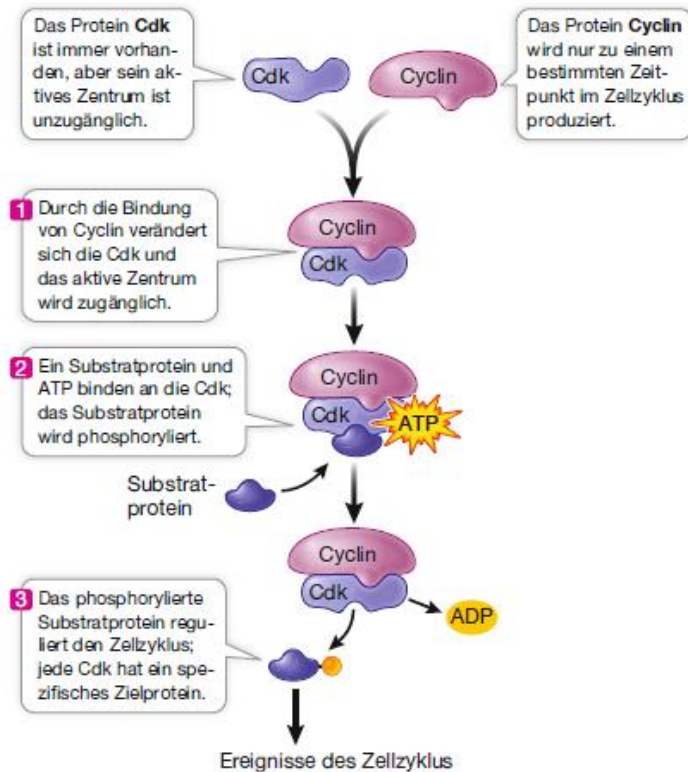
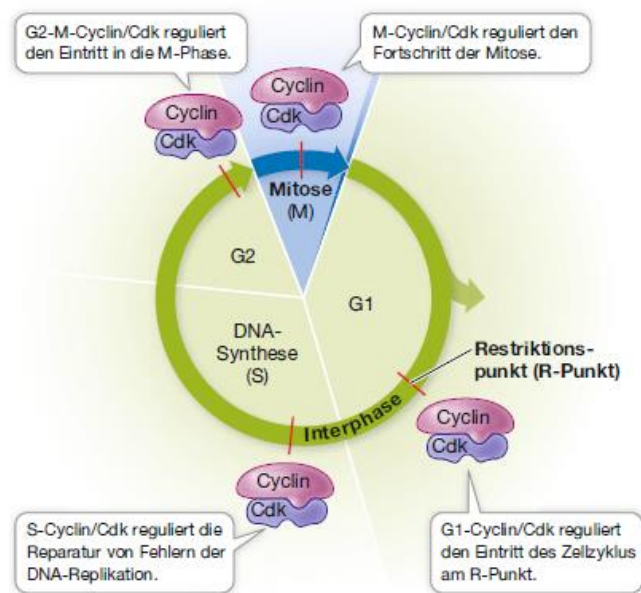


Abbildung 131 Cdk und Zellzyklus

Die verschiedenen Cyclin/Cdk-Komplexe fungieren als Kontrollpunkte des Zellzyklus in Form von Signalwegen, die den Fortschritt des Zellzyklus regulieren. Ist beispielsweise eine Zelle durch Strahlung oder Gift grundlegend geschädigt, kann sie daran gehindert werden, den Zellzyklus abzuschließen.

Als Beispiel soll hier der G1-Kontrollpunkt (R) dienen. Wird die DNA während der G1-Phase geschädigt, führt ein Signalweg zur Produktion des Proteins p21. Das p21-

Protein bindet an die G1-S-Cdk, sodass sie Cyclin nicht mehr binden kann. Dadurch bleibt die Cdk inaktiv und der Zellzyklus hält an, während die DNA repariert wird.

Wenn der Signalweg für DNA-Schädigung nicht mehr aktiv ist, wird p21 abgebaut, sodass die Funktion des Cyclin/Cdk-Komplexes zum Tragen kommt und sich der Zellzyklus fortsetzt. Wenn die DNA-Schädigung gravierend ist und nicht repariert werden kann, durchläuft die Zelle den programmierten Zelltod (die Apoptose). Solche Kontrollen verhindern, dass sich geschädigte Zellen vermehren und einem Organismus schaden können.

Die Cyclin/Cdk-Komplexe geben der Zelle die Möglichkeit, den Fortschritt des Zellzyklus intern zu regulieren. Der Zellzyklus wird jedoch auch durch äußere Signale beeinflusst. Nicht alle Zellen in einem Organismus durchlaufen den Zellzyklus gleich schnell. Manche Zellen verlassen ihn und treten in die G0-Phase ein, oder sie durchlaufen ihn nur sehr langsam und teilen sich entsprechend selten. Wenn sich solche Zellen teilen sollen, müssen sie durch externe chemische Signale stimuliert werden, die man als Wachstumsfaktoren bezeichnet. Z. B. wenn wir uns schneiden und bluten sammeln sich spezialisierte Zellfragmente, die Blutplättchen (Thrombocyten), an der Wunde und setzen die Blutgerinnung in Gang. Die Thrombocyten produzieren das Protein PDGF (plateletderived growth factor), das zu den angrenzenden Zellen in der Haut diffundiert und sie zur Zellteilung und Wundheilung anregt. Wachstumsfaktoren binden an spezifische Rezeptoren auf ihren Zielzellen und aktivieren Signaltransduktionswege, an deren Ende die Synthese von Cyclin steht, wodurch die Cdks und der Zellzyklus aktiviert werden.

Mitose

Während der S-Phase wird die DNA bzw. Chromosomen verdoppelt, es entstehen die Schwesterchromatiden. Diese werden während der G2-Phase über den größten Teil ihrer Länge von einem Proteinkomplex zusammengehalten, den man als Cohesin bezeichnet. Am Ende der G2-Phase und zu Beginn der Mitose umhüllen die sogenannten Kondensine, eine andere Gruppe von Proteinen, die DNA-Moleküle, wodurch diese durch Faltung kompakter werden. Außerdem wird in der G2-Phase der Beginn der Bildung des Spindelapparat initiiert. Dieser gehört zum Cytoskelett und trennt die Schwesterchromatiden bei der Mitose voneinander. Die Mitose findet in verschiedenen Schritten statt (Abb. 132).

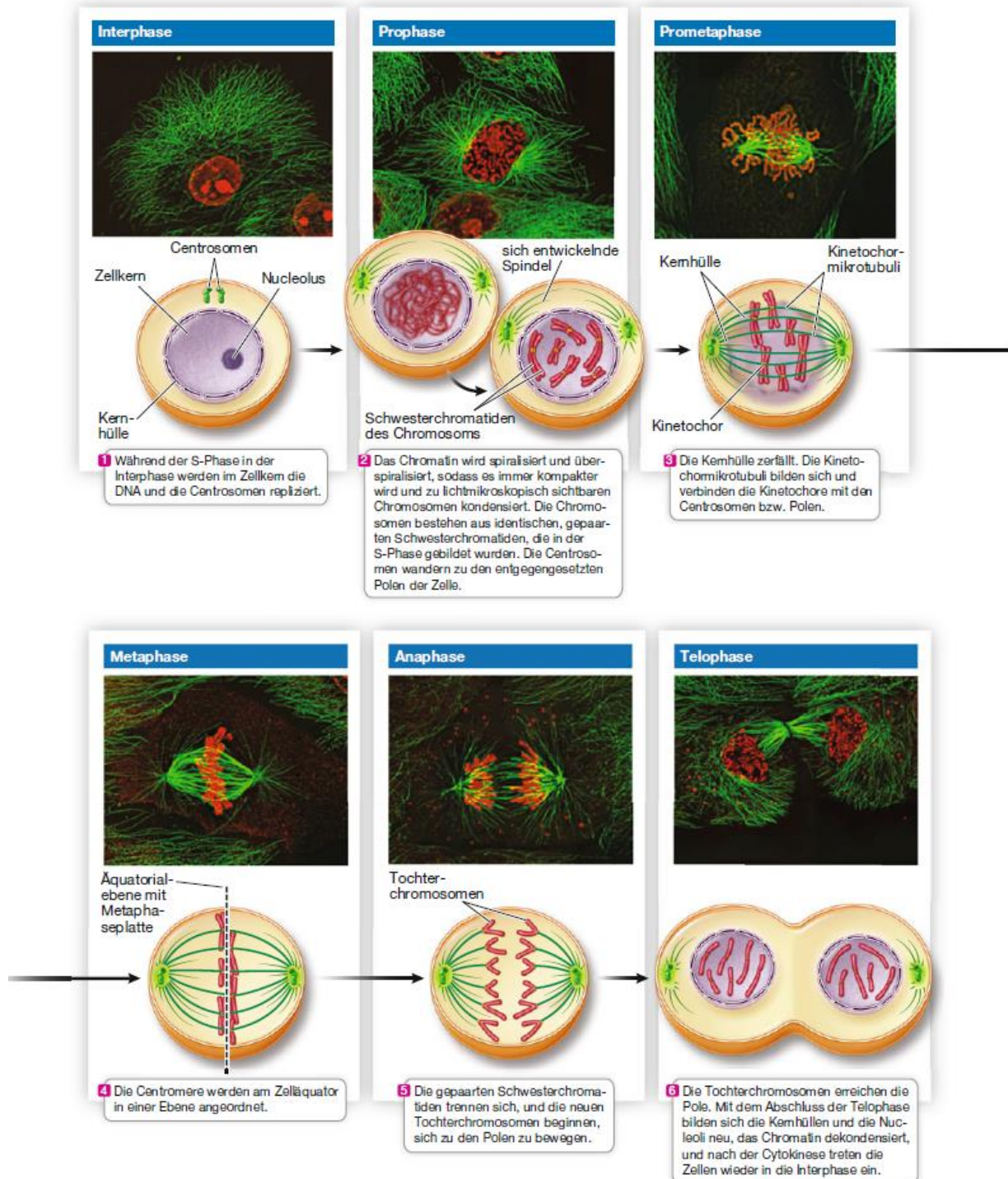


Abbildung 132 Mitose

1. Interphase: G1 (Wachstum; am Ende liegt der Restriktionspunkt), S (DNA-Replikation), G2 (Bildung des Spindelapparats beginnt)
2. Prophase: Kondensation der DNA der Chromosomen, die lichtmikroskopisch sichtbar werden; Zusammenbau des Spindelapparates aus den Centrosomen
3. Prometaphase: Zellkernhülle löst sich auf; Chromosomen heften sich an die Spindel

4. Metaphase: Anordnung der Chromosomen zu einer Äquatorialplatte (Metaphasenplatte)
5. Anaphase: Trennung der Chromatiden; Wanderung zu den Polen der Zelle
6. Telophase: Dekondensieren der Tochterchromosomen; Bildung der neuen Kernhüllen
7. Cytokinese

Nach ihrer Trennung bezeichnet man die Schwesterchromatiden auch als Tochterchromosomen. Zur besseren Unterscheidung spricht man bis zur Trennung auch von Zwei-Chromatid-Chromosomen und nach der Trennung von Ein-Chromatid-Chromosomen.

Spindelapparat

Der Spindelapparat ist eine dynamische Struktur aus Mikrotubuli, die die Schwesterchromatiden bei der Mitose voneinander wegbewegt. Die Mikrotubuli sind die größten Bestandteile des Cytoskeletts (Abb. 133). Sie sind zwischen 15 und 25 nm groß. Zusammengesetzt sind die Mikrotubuli aus dem runden Protein namens Tubulin. Ihr Aufbau ist durch eine Basisstruktur, das Protofilament, gegeben.

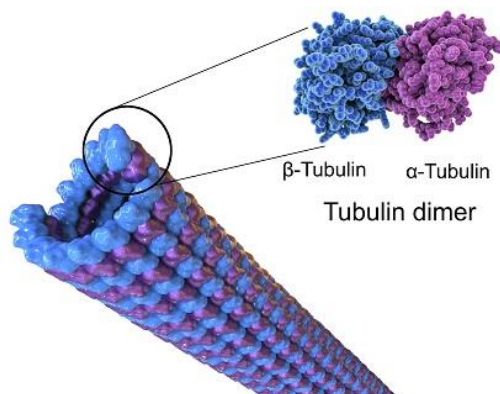


Abbildung 133 Mikrotubuli

Die Protofilamente sind dabei aus den sogenannten alpha-Tubulinen an einem Ende und beta-Tubulinen am anderen Ende aufgebaut. Das sind kugelförmige Proteine aus ca. 450 Aminosäuren. Ein Mikrotubulus besteht aus ungefähr 13 dieser Protofilamente. Die Protofilamente sind dabei durch das sogenannte Kopf-Schwanz-Prinzip aufgebaut. Das bedeutet, dass ein alpha-Tubulin (Kopf) immer mit einem beta-Tubulin verbunden ist.

Diese Art der Verknüpfung tritt in sehr vielen Verbindungen auf. Analog dazu existieren auch Kopf-Kopf-Verknüpfungen und Schwanz-Schwanz-Verknüpfungen.

Über das alpha-Tubulin sind die Mikrotubuli an das sogenannte Mikrotubulus-Organisationszentrum (MTOC) gebunden. Dies ist quasi der Startpunkt des Wachstums, also der Verknüpfung der Mikrotubuli. Die Mikrotubuli sind zu einem Großteil für den Stofftransport innerhalb des Zytoskeletts verantwortlich, mit dem wir uns später befassen werden. Seine zweite Aufgabe, die Ausbildung des Spindelapparates, wird hier besprochen.

Bevor sich der Spindelapparat bilden kann, wird seine Orientierung durch das Centrosom (Zentralkörperchen) festgelegt. Dies ist ein membranloses Organell im Cytoplasma in der Nähe des Zellkerns. Bei den Tieren enthält das Centrosom ein Paar von Centriolen, die jeweils die Form eines kurzen Zylinders besitzen, der aus neun Dreiergruppen von Mikrotubuli besteht. Während der S-Phase verdoppelt sich das Centrosom, und zu Beginn der Prophase trennen sich die beiden Centrosomen voneinander und wandern zu den entgegengesetzten Enden der Kernhülle. Sie bestimmen so die beiden Zellpole, auf die sich die Chromatiden während der Anaphase zu bewegen. Die Zellen von Pflanzen und Pilzen besitzen keine Centriolen, sondern bilden ein Mikrotubuliorganisationszentrum (MTOC) an jedem Zellpol, das die gleiche Funktion erfüllt. Die Positionen der beiden Centrosomen legen die Ebene fest, an der sich eine tierische Zelle teilt. Dadurch bestimmen sie die räumliche Beziehung zwischen den beiden neuen Zellen.

Während der Prophase wird der entsprechende Spindelapparat aufgebaut. In diesem Stadium verschwindet der größte Teil des Cohesins, das seit der S-Phase die beiden DNA-Moleküle zusammengehalten hat.

Der Spindelapparat wird aus drei Mikrotubulus-Arten aufgebaut. Das sind die Polar-, Astral- und Kinetochor-Mikrotubuli (Abb. 134).

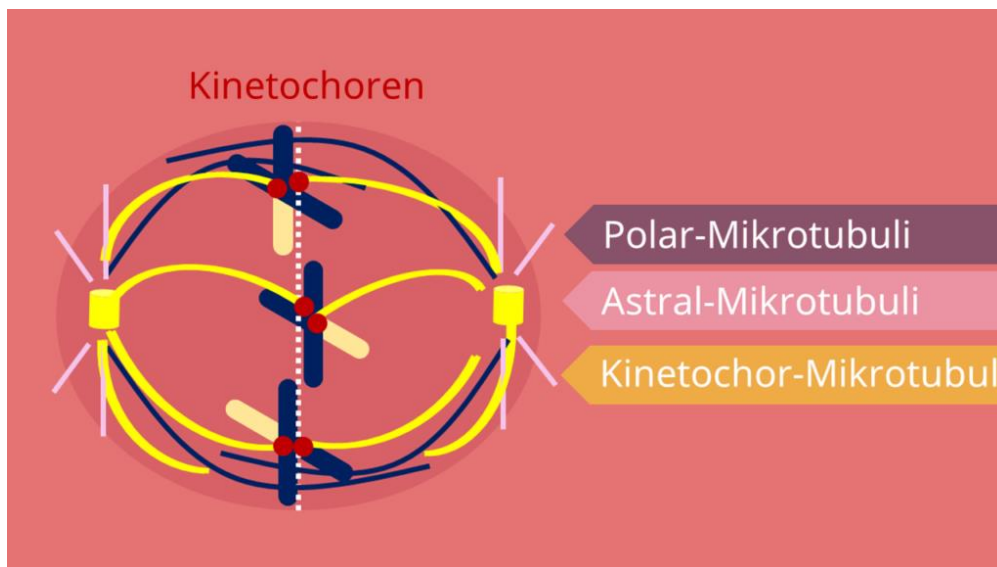


Abbildung 134 Spindelapparat

Die polaren Mikrotubuli sind von Zellpol zu Zellpol aufgespannt. Sie können somit über die Äquatorialebene des Spindelapparates hinausreichen.

Die astralen Mikrotubuli stellen die Verbindung zum Cytoskelett her. Sie sind dabei sternförmig um die Zellpole angeordnet.

Die Kinetochor-Mikrotubuli sind wahrscheinlich die wichtigsten der drei Arten. Sie sind maßgeblich an der Trennung der Chromosomen beteiligt. Sie werden in der Prophase der Mitose gebildet. In der darauffolgenden Prometaphase binden sie an die sogenannten Kinetochoren der Chromosomen. Diese sind kleine Proteinstrukturen, die am Centromer der Chromosomen sitzen. Sie sind dabei eine Art Bindungsstelle für

einen Kinetochor-Mikrotubulus. Nachdem an den Kinetochor angedockt wurde, werden die Chromosomen auseinandergezogen.

Cytokinese

Nach der Mitose, also nach Bildung der beiden Tochterkerne, teilt sich das Cytoplasma der Zelle im Rahmen der Cytokinese. Bei diesem Vorgang bestehen zwischen Tier- und Pflanzenzellen zwei grundlegende Unterschiede. Bei Tierzellen beginnt die Cytokinese mit einer Einschnürung der Plasmamembran, als ob sich zwischen den beiden Zellkernen ein unsichtbarer Gummiring bildet und zusammenzieht. Er wird von Actinfilamenten und dem assoziierten Motorprotein Myosin an der cytoplasmatischen Oberfläche der Plasmamembran gebildet. Actinfilamente sind Bestandteile des Cytoskeletts. Motorproteine sind intrazelluläre Transportproteine, die unter Energieverbrauch Moleküle in der Zelle transportieren können und mit dem Cytoskelett verbunden sind.

Diese beiden Proteine interagieren und bewirken dadurch eine Kontraktion, sodass durch die Einschnürung schließlich zwei Zellen entstehen. Actinfilamente bilden sich sehr schnell aus Actinmonomeren, die bereits in der Interphase als Bausteine des Cytoskeletts vorhanden sind.

Das Cytoplasma von Pflanzenzellen teilt sich auf andere Weise, da Pflanzen feste Zellwände besitzen. In Pflanzenzellen treten während des Abbaus der Spindel nach der Mitose entlang der Zellteilungsebene, etwa in der Mitte zwischen den beiden Tochterkernen, vom Golgi-Apparat abgeschnürte Membranvesikel auf. Die Vesikel werden durch das Motorprotein Kinesin entlang von Mikrotubuli (ebenso Bestandteil des Cytoskeletts) bewegt und verschmelzen mit ihresgleichen zur Bildung der beiden trennenden Plasmamembranen. Gleichzeitig bildet sich aus ihren Inhaltsstoffen eine Zellplatte, das Anfangsstadium der neuen trennenden Zellwand.

Nach der Cytokinese enthält jede Tochterzelle alle Bestandteile einer vollständigen Zelle. Durch die Mitose wird die exakte Verteilung der Chromatiden sichergestellt. Im Gegensatz dazu werden Organellen wie Ribosomen, Mitochondrien und Chloroplasten nicht unbedingt gleichmäßig zwischen den Tochterzellen aufgeteilt, es müssen nur in jeder Zelle genügend davon vorhanden sein. Dementsprechend gibt es für die Organellen keinen so präzisen und offensichtlichen Verteilungsmechanismus wie für die Mitose. Umso erstaunlicher ist, dass die Verteilung der Organellen durchweg funktioniert, und es ist weitgehend unverstanden, wie sie gesteuert wird.

Kapitel 14: Meiose, Rekombination und Kernphasenwechsel

Im letzten Beitrag befassten wir uns mit der Mitose, bei dem das genetische Material gleichmäßig auf die Tochterzellen weitergegeben wird. Wir erinnern uns: Menschen haben 46 Chromosomen: 44 Autosomen und die zwei Gonosomen, auch Geschlechtschromosomen genannt (Bei Frauen XX und beim Mann XY). Genauer gesagt haben wir 23 homologe Chromosomenpaare, d. h. jedes Chromosom kommt doppelt vor. Man spricht von diploid. Die Ausnahme stellt natürlich das Y-Chromosom des Mannes dar, welches nicht homolog zum X-Chromosom ist; dennoch bilden beide das Paar der Geschlechtschromosomen. Der diploide Chromosomensatz kommt zustande, weil bei der sexuellen Vermehrung Eizelle und Spermazelle miteinander verschmelzen, ein Vorgang den man als Befruchtung bezeichnet. Eizelle und Spermazelle sind die Keimzellen und diese haben nur einen halben Chromosomensatz, man spricht von haploid. Wie dieser halbe Chromosomensatz entsteht, wird in diesem Beitrag geklärt.

Meiose

Die Meiose lässt sich in Meiose I und Meiose II gliedern (Abb. 135).

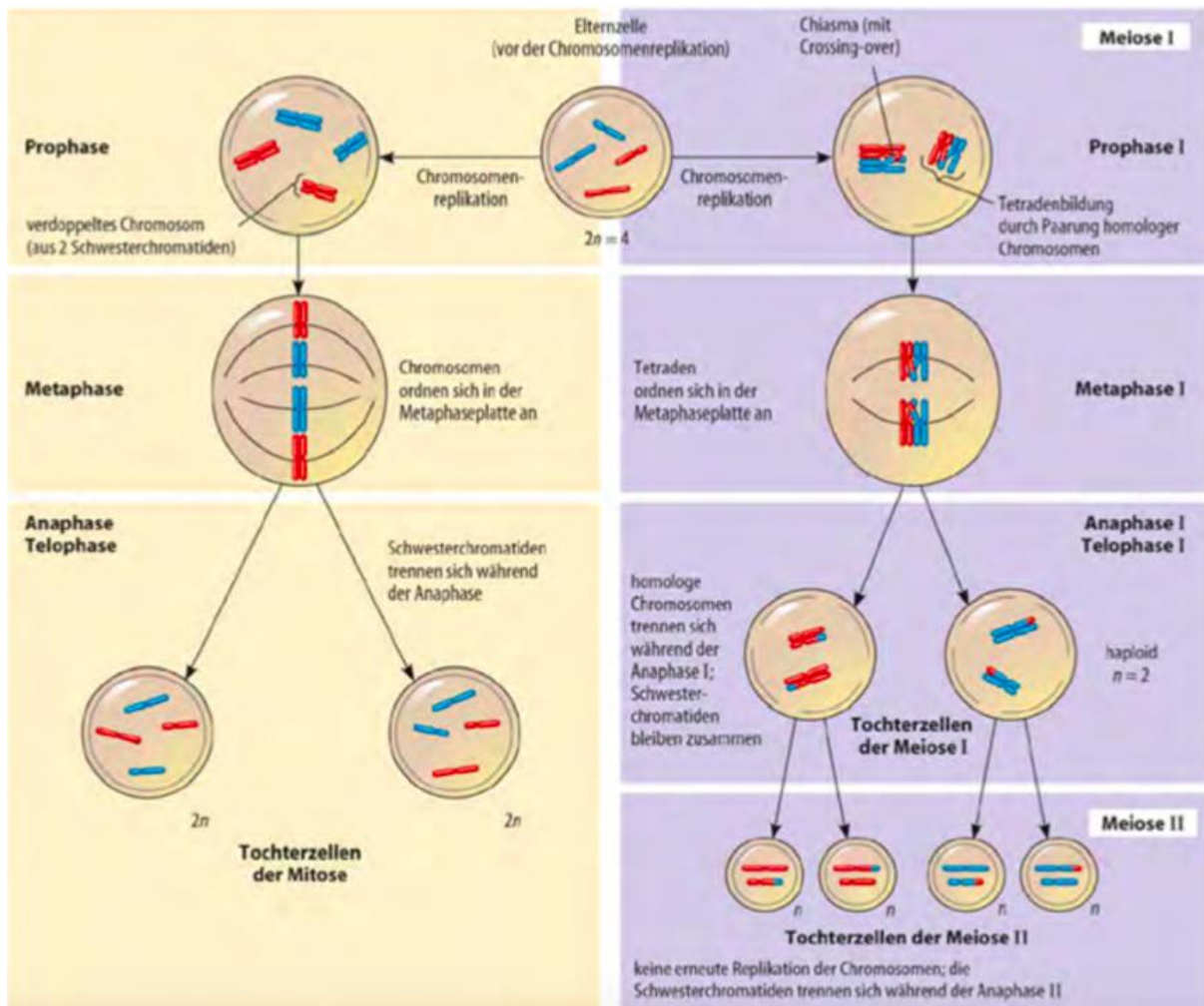


Abbildung 135 Unterschied Mitose und Meiose

Der Meiose I geht wie der Mitose eine Interphase mit einer S-Phase voraus, in der jedes Chromosom repliziert wird. Das führt dazu, dass jedes Chromosom aus zwei Schwesterchromatiden besteht, die von Cohesinen zusammengehalten werden.

Die Meiose I beginnt mit einer langen Prophase I, bei der sich die Chromosomen deutlich verändern. Die homologen Chromosomen lagern sich Seite an Seite zu Paaren zusammen, wodurch Tetraden entstehen, also Gruppen aus jeweils vier Chromatiden. Bei der Mitose tritt eine solche Tetradenbildung nicht auf. Die vier Chromatiden stammen von den beiden Partnern eines jeden homologen Chromosomenpaares.

Bei der Mitose verhält sich jedes Chromosom unabhängig von seinem homologen Gegenstück, und seine beiden Schwesterchromatiden werden in der Anaphase zu den beiden entgegengesetzten Polen gebracht. Jeder Tochterzellkern enthält schließlich wieder $2n$ Chromosomen.

Die Meiose verläuft hier deutlich anders.

In der Meiose I paaren sich in der Synapse die Chromosomen mütterlichen Ursprungs mit ihren väterlichen Gegenstücken. Da jedes Chromosom aufgrund der DNA-Replikation aus zwei identischen Schwesterchromatiden besteht, liegen jetzt Tetraden

aus vier sich entsprechenden Chromatiden vor; da die beiden Chromosomen nicht identisch sind, liegen in der Mitte der Tetrade zwei Nicht-Schwesterchromatiden aneinander. Die Paarung mütterlicher und väterlicher homologer Chromosomen tritt bei der Mitose nicht auf.

In der Metaphase I ordnen sich die homologen Chromosomenpaare an der Äquatorialplatte an und werden in der Anaphase I zu den jeweiligen Polen gezogen.

Die Tetraden trennen sich wieder in einzelne Chromosomen, die jeweils aus zwei Schwesterchromatiden bestehen.

Am Ende der Meiose I, genauer der Telophase I, bilden sich zwei Zellkerne; jeder enthält die Hälfte der ursprünglichen Chromosomen (eines von jedem homologen Paar). Da sich die Centromere bis dahin nicht getrennt haben, bestehen diese Chromosomen immer noch aus den beiden Schwesterchromatiden. Die Schwesterchromatiden werden in der Meiose II getrennt, vor der keine DNA-Replikation stattfindet.

Meiose II verläuft im Prinzip wie die Mitose:

In Prophase II kondensieren die Chromosomen, in der Metaphase II ordnen sich die Chromosomen an die Äquatorialplatte und werden in Anaphase II zu den Polen gezogen. Die Schwesterchromatiden trennen sich schließlich und werden eigenständige Tochterchromosomen.

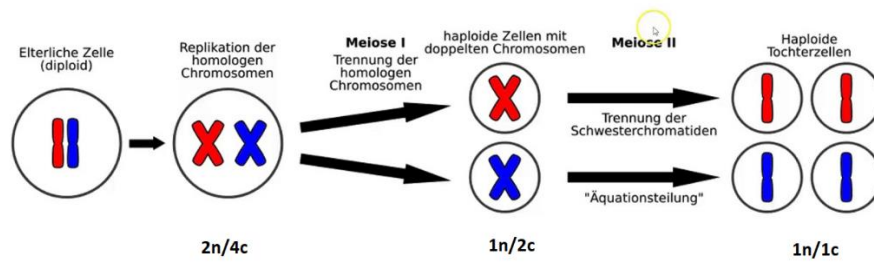
Am Ende der Meiose haben wir also 4 Keimzellen mit dem halben Chromosomensatz. Aber diese vier Zellen sind genetisch nicht identisch.

Um es zusammenzufassen (Abb. 136):

Die Elternzelle ist diploid, das Genom wird repliziert, es entstehen Tetraden (oder anders gesagt: $2n/4c$). In der Meiose I ordnen sich die homologen Chromosomenpaare paarweise in der Äquatorialebene an und werden voneinander getrennt. Es entstehen zwei haploide Zellen. Die einzelnen Chromosomen bestehen jedoch noch aus je zwei Chromatiden ($1n/2c$). In der Meiose II werden die einzelnen Schwesterchromatiden geteilt und es entstehen 4 haploide Zellen. Jedes Chromosom besteht nur noch aus einem Chromatid ($1n/1c$).

Weil bei der Meiose der Chromosomensatz halbiert wird, spricht man auch von einer Reduktionsteilung.

Allgemeiner Ablauf der Meiose



BRUNNEN UND
SCHÖNHAUT

Abbildung 136 Meiose

Typen von Keimzellen

Bei einer Meiose entstehen aus einer diploiden Zelle vier haploide. Doch wenn man sich die haploiden Keimzellen, auch Gameten genannt, näher anschaut, sehen sie nicht immer gleich aus. Die beiden sich befruchtenden Gameten gehören zwei verschiedenen, als Geschlechter (Sex) bezeichneten Typen an. Entsprechend gibt es auch zwei unterschiedliche Gametentypen. Die kleineren bzw. beweglichen Gameten sind die männlichen Gameten, die größeren bzw. unbeweglichen Gameten bezeichnet man als weiblich. Die Organismen, die in der Lage sind entsprechende Gameten zu bilden werden als männlich oder weiblich eingestuft. Bei der Befruchtung kann stets nur ein Gamet des einen Typen mit dem anderen verschmelzen. Die Gameten sind also die Ursache der bipolaren Zweigeschlechtlichkeit, was sich durch ihre funktionelle Bedeutung wie auch ihre chemische Struktur, Zelloberfläche etc. ergibt.

Sehen die Keimzellen morphologisch (aber nicht zwingend chemisch oder funktionell!) ähnlich, spricht man von einer Isogamie. Im Falle einer solchen Isogamie werden die beiden, nur an ihren funktionellen Eigenschaften unterscheidbaren Sorten willkürlich als + und – bezeichnet, sofern nicht ein Verhaltensunterschied bei der Einleitung der Befruchtung eine Identifizierung des Geschlechtes erlaubt. Lassen sich die Gameten anhand ihres Aussehens unterscheiden, spricht man von Anisogamie.

Sind beide Gameten begeißelt unterscheidet man zwischen männlichen Mikrogameten und weiblichen Makrogameten.

Ist nur ein Gametentyp begeißelt und damit beweglich, wird dieser – unabhängig von seiner Größe – als männlich definiert.

Sind beide Eigenschaften kombiniert, so sind die kleinen und beweglichen Gameten die männlichen Spermatozoen (bei Tieren Spermien genannt) und die großen unbeweglichen Gameten sind die weiblichen Eier. Diese Kombination bezeichnet man auch als Oogamie. Sie kommt bei allen Tieren (also auch uns Menschen) vor.

Einen Überblick über die verschiedenen Keimzellentypen gibt Abb. 137:

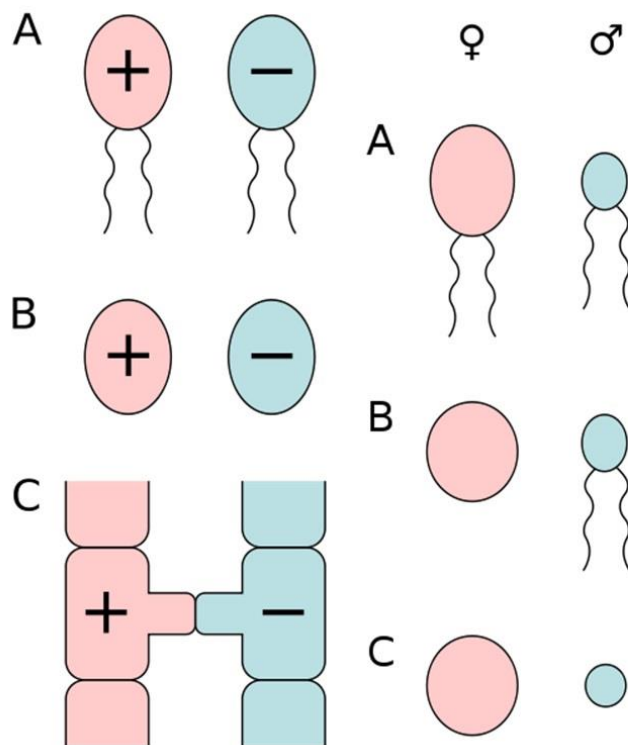


Abbildung 137 die unterschiedlichen Gametentypen. Linke Spalte: A) aktiv bewegliche Isogameten, B) unbewegliche Isogameten, C) Gametangiogamie: Befruchtung ohne Gameten bei den Jochpilzen. Rechte Spalte: A) aktiv bewegliche Anisogameten, B) Eizelle und Spermium (Oogamie), C) unbewegliche Anisogameten (Eizelle und Spermium).

Der Grund weshalb Eizellen übrigens größer sind als Spermien, liegt in einer Sonderfunktion der Meiose bei Eizellen. Hier entstehen zwar auch aus einer diploiden Mutterzelle 4 haploide Zellen. Aber drei dieser haploiden Zellen werden zu sogenannten Polkörpern. Die Teilung der 4 Zellen ist inäqual. Von einer großen Zelle wird nur eine kleine Zelle abgegliedert, die sich dann entweder auflöst oder nochmal teilt. So entsteht eine sehr große Eizelle und 3 sogenannte Polkörper, die für die Befruchtung funktionslos sind (Abb. 138).

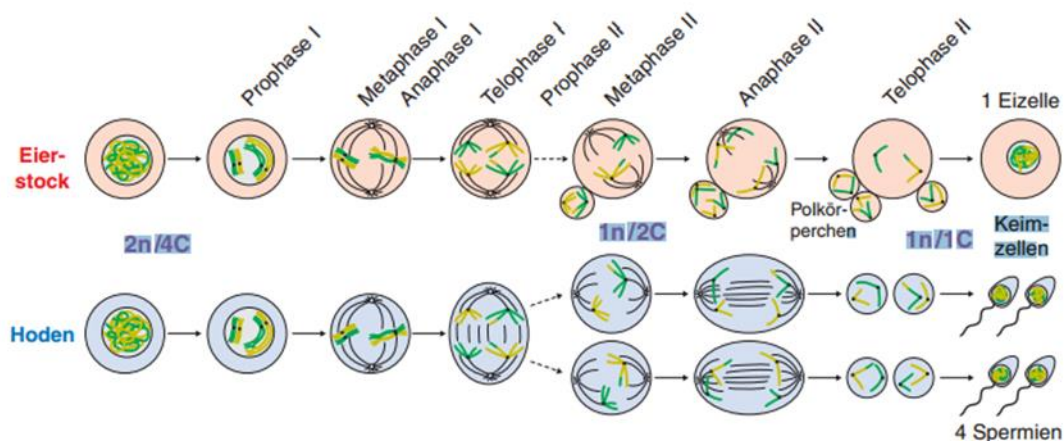


Abbildung 138 Spermien- und Eizellenproduktion

Die Produktion unterschiedlicher Gameten geht natürlich einher mit der Ausbildung unterschiedlicher Strukturen ihrer Erzeuger: Das gilt sowohl für primäre Geschlechtsorgane, sowie die sekundären. Im Extremfall bildet sich ein regelrechter Sexualdimorphismus aus, bei dem Männchen und Weibchen so unterschiedlich aussehen, als ob sie anderen Arten angehören könnten (vgl. ALBERTS et al. 2017, CZIHAK 1981, WOLPERT et al. 2007). Da in der Natur also auch immer zwei unterschiedliche Gametentypen pro Art ausgebildet werden, ist auch hier die Einteilung in zwei Geschlechter sinnvoll. Ein drittes Geschlecht würde biologisch nur dann Sinn machen, wenn es triploide Arten gibt, die aus drei unterschiedlichen Keimzellen hervorgehen. Natürlich und das muss gesagt werden: Aus der Bildung der Keimzellen lassen sich zwar biologische Geschlechter definieren, das sagt aber nicht ansatzweise etwas über Geschlechterrollen und gesellschaftliche Fragen etwas aus.

Crossing Over

Die Meiose und sexuelle Fortpflanzung ermöglichen des Weiteren die Erzeugung genetischer Vielfalt. Während die Mitose, salopp formuliert, eine Form des Klonens ist, sorgt die Meiose für genetische Rekombination.

Während der Prophase I und der Metaphase I setzt das Chromatin seine Spiralisierung und Komprimierung fort, sodass die Chromosomen immer dicker erscheinen. Zu einem bestimmten Zeitpunkt werden die homologen Chromosomen durch einen Spindelapparat aus Mikrotubuli voneinander weggezogen, vor allem in der Nähe der Centromere, aber sie bleiben über die Cohesine noch miteinander verbunden. Später in der Prophase I nehmen die Bereiche, in denen diese Anheftungen auftreten, die Form eines X an und man bezeichnet sie als Chiasmata. Ein Chiasma macht den Austausch von genetischem Material zwischen Nicht-Schwesterchromatiden auf homologen Chromosomen sichtbar – in der Genetik spricht man hier von einem Crossing-over (Abb. 139).

Ein Crossing-over führt zu rekombinanten (neu kombinierten) Chromatiden. Diese Rekombination des Erbmaterials erhöht die genetische Variabilität unter den Produkten der Meiose, indem genetische Information innerhalb der jeweiligen homologen Paare verschoben wird.

Das Crossing-over ist nur eine der Ursachen für die genetische Vielfalt der Meioseprodukte. Die andere Ursache ist die unabhängige Verteilung. Dabei erhält jede haploide Zelle aus der diploiden Zelle einen vollständigen Satz von Genen, aber nur jeweils ein Gen von jedem Genpaar von jedem Elternteil. Die Entscheidung, welches Chromosom aus einem homologen Paar in der Anaphase I in welche Tochterzelle gelangt, erfolgt rein zufällig. Stellt euch beispielsweise vor, in einem diploiden Zellkern seien zwei homologe Chromosomenpaare vorhanden.

Ein bestimmter Tochterzellkern könnte das väterliche Chromosom 1 und das mütterliche Chromosom 2 erhalten. Oder er könnte das väterliche Chromosom 2 und das mütterliche Chromosom 1 oder beide mütterlichen oder beide väterlichen

Chromosomen erhalten. Das alles hängt davon ab, wie sich die homologen Paare in der Metaphase I zufällig anordnen. Je größer die Anzahl der Chromosomen, desto geringer ist die Wahrscheinlichkeit, dass die ursprünglichen elterlichen Kombinationen erneut entstehen, und umso größer ist das Potenzial für das Entstehen von genetischer Vielfalt.

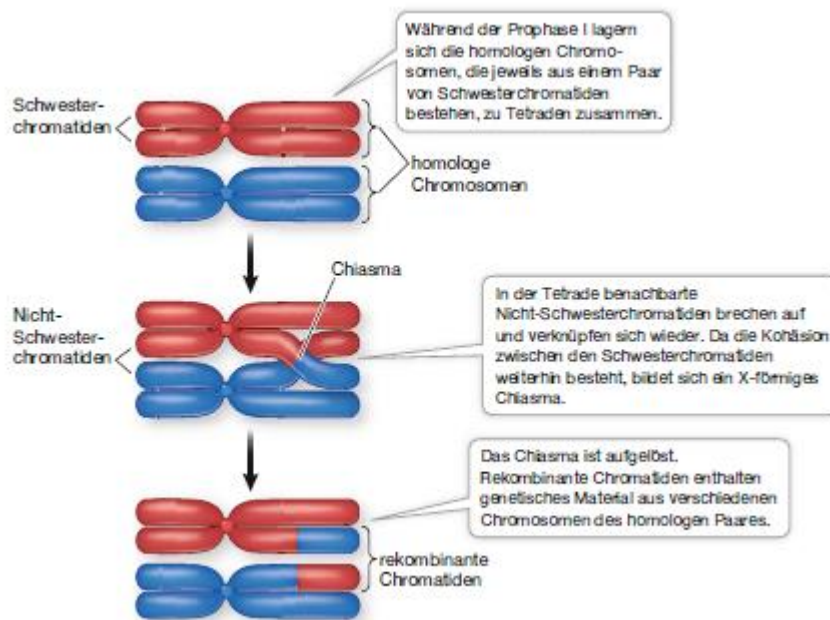


Abbildung 139 Crossing over

Kernphasenwechsel

Dieser regelmäßige Wechsel zwischen einem durch die Meiose eingeleiteten haploiden Zustand (Haplophase) und einem aus der Befruchtung resultierenden diploiden Zustand (Diplophase) ist für alle Eukaryoten, die sich sexuell Vermehren, kennzeichnend. Organismen haben dabei also immer einen Wechsel zwischen einer haploiden Phase (Gameten) und einer diploiden. Je nachdem welche der beiden Phasen überwiegt, kann man sie drei verschiedenen Typen zuordnen (Abb. 140)

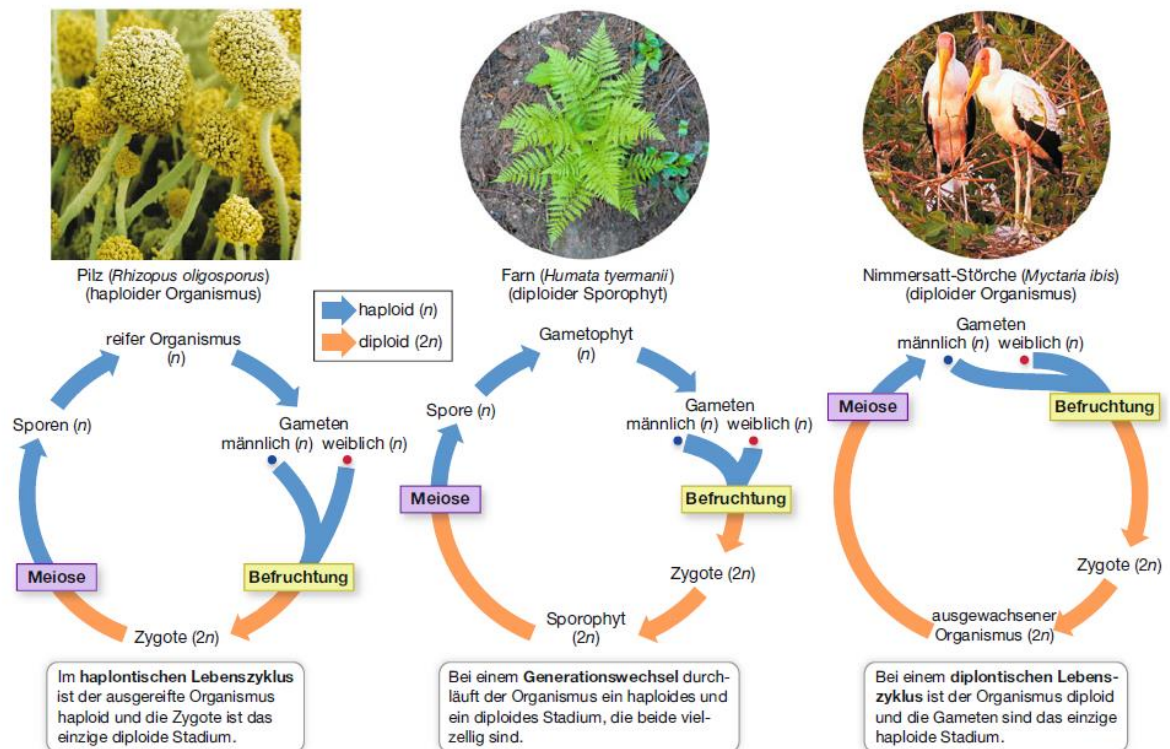


Abbildung 140 Kernphasenwechsel

Haplonten

Bei Haplonten spielt sich fast der gesamte Entwicklungszyklus in der Haplophase ($1n$) ab. Nur die Zygote (= befruchtete Eizelle) besitzt einen diploiden Chromosomensatz ($2n$). Diese entsteht durch die Verschmelzung zweier Gameten. Die so entstandene Zygote vollzieht danach eine Reduktionsteilung (Meiose). So entstehen wieder 4 haploide Zellen (Gonen) und aus den haploiden Zellen kann durch Mitose wieder ein haploider Vielzeller entstehen. Zu den Haplonten gehören viele Pilze, viele Algen und einige tierische Einzeller, wie die Flagellaten.

Diplonten

Die Diplonten sind quasi das Gegenteil der Haplonten. Bei ihnen vollzieht sich die Lebensphase fast ausschließlich in der Diplophase. Nach der Befruchtung erfolgt, anders als bei den Haplonten, keine Reduktionsteilung (Meiose), sondern die Zygote entwickelt sich zum diploiden Organismus. Einige dieser Zellen bilden dann die haploiden Gameten, wodurch sich die Organismen wieder sexuell vermehren können. Neben einigen Algen und Pilzen gehören alle vielzelligen Tiere (Metazoa) und die Ciliata (Wimpertierchen, Einzeller) zu den Diplonten.

Diplo-Haplonten

Die Diplo-Haplonten (auch Haplo-Diplonten genannt) haben eine Zwischenstellung zu den zwei oben genannten Typen. Es sind Lebewesen, bei deren Fortpflanzung abwechselnd haploide und diploide Generationen auftreten. Hierzu gehören die meisten Pflanzen. Kennzeichnend für Diplo-Haplonten ist, dass sowohl in der haploiden wie in der diploiden Phase Mitosen erfolgen. Die haploide Generation pflanzt sich geschlechtlich fort, indem sie weibliche und männliche Gameten bildet. Sie wird deshalb als Gametophyt bezeichnet. Durch die Vereinigung zweier Gameten verschiedenen Geschlechts und Befruchtung entsteht eine diploide Zygote. Die Zygote wächst zum sogenannten Sporophyten. Die Sporophyten bilden nach einer Reduktionsteilung (Meiose) haploide Sporen. Aus diesen haploiden Sporen wächst der haploide Gametophyt heran und der Zyklus beginnt von vorne. In der Evolutiongeschichte der Pflanzen kommt es dabei zur Reduktion des haploiden Gametophyten. Bei den Moosen ist der Gametophyt die dominante Phase (nämlich die grüne Moospflanze). Schon bei den Farnen ist der diploide Sporophyt dominant und bei den Samenpflanzen ist die eigentliche Pflanze der diploide Sporophyt, während der haploide Gametophyt auf wenige Zellen reduziert ist (männlicher Pollenschlauch, bestehend aus drei Zellen und weiblicher Embryosack in der Samenanlage, bestehend aus 7-8 Zellen).

Kapitel 15: Endoplasmatisches Reticulum, Golgi Apparat, Lysosom

Wir haben verschiedene Zellorganellen schon kennengelernt, allen voran Mitochondrien und Chloroplasten bei Zellatmung und Photosynthese, den Zellkern bei der DNA und die Grundlagen der Biomembranen.

In diesem Beitrag wenden wir uns weiteren Zellorganellen zu, vornehmlich dem endoplasmatischen Reticulum und dem Golgi-Apparat.

Endoplasmatisches Reticulum (ER) und Golgi-Apparat (GA)

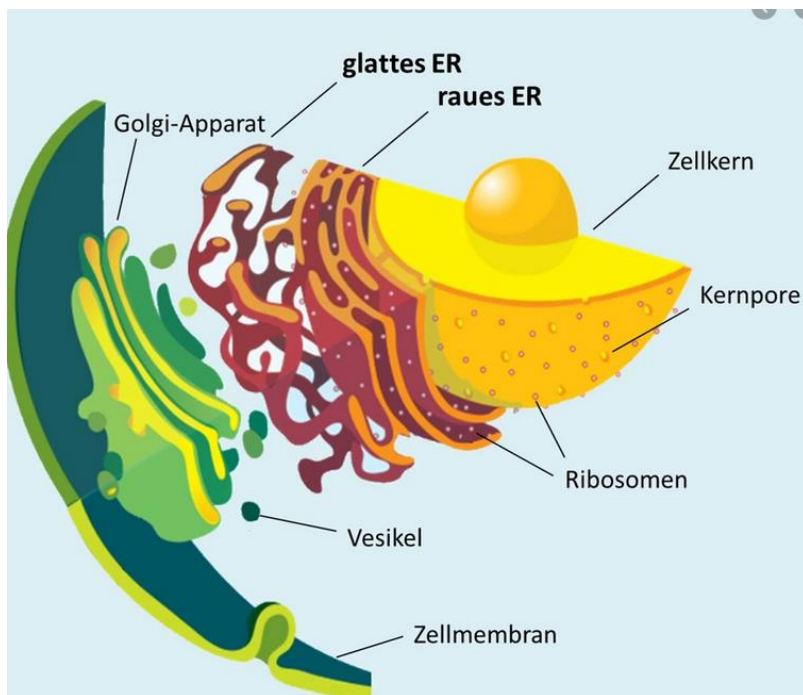


Abbildung 141 Endoplasmatisches Reticulum (ER) und Golgi-Apparat (GA)

Das endoplasmatische Reticulum liegt in zwei Formen vor, dem glatten ER (SER) und dem rauen ER (RER). Auf dem rauen ER befinden sich die Ribosomen, kleine Zellorganellen, die bei der Produktion von Proteinen eine wichtige Rolle spielen.

Das endoplasmatische Reticulum erfüllt mehrere wichtige Aufgaben. Grundsätzlich ist es für die Signalübertragung innerhalb der Zelle verantwortlich. Diese Signalübertragung kann man sich als eine Aufnahme von Calcium-Ionen vorstellen. Das ist insbesondere bei Muskelzellen und Nervenzellen sehr wichtig.

Das glatte ER ist vor allem bei Stoffwechselvorgängen wichtig. Enzyme, die sich innerhalb des glatten ER befinden, können Lipide herstellen.

Weil sich am rauen ER die Ribosomen befinden, ist dieses vor allem bei der Herstellung von Proteinen beteiligt. Eine weitere Hauptfunktion des rauen ER ist die Bildung einer Kernmembran.

Das sogenannte Sarkoplasmatische Retikulum (SR) ist eine spezielle Form des endoplasmatischen Retikulums. Es befindet sich insbesondere in der glatten und der quergestreiften Muskulatur, besonders in den Muskelfasern. In der Nähe des SR befinden sich auch viele Mitochondrien.

Die Hauptaufgabe des SR ist die Speicherung von Calcium-Ionen. Wenn ein Nervensignal an die Muskeln geliefert wird, schüttet das SR die Calcium-Ionen in das Plasma der Muskelzellen aus. Dadurch sind die Calcium-Ionen in der Lage, ein Zusammenziehen des Muskels (Kontraktion) einzuleiten. Somit ist das SR ein wichtiger Vermittler, um jede deiner Bewegungen ausführen zu können.

Ein weiteres wichtiges internes Membransystem stellt der Golgi-Apparat dar (siehe Abb. 1). Eine seiner Aufgaben ist es, Proteine zu verarbeiten und umzuwandeln. Die Proteine kommen dabei vor allem vom endoplasmatischen Retikulum (ER). Außerdem ist er in der Lage, diese Proteine in Vesikeln zu binden. Diese können die Proteine zu verschiedensten Bereichen innerhalb der Zelle transportieren.

Der Golgi-Apparat besteht aus mehreren gestapelten Membran-Zisternen, die auch als Dictyosomen bezeichnet werden. Die Gesamtheit der Dictyosomen innerhalb der Zelle wird als Golgi-Apparat bezeichnet.

Um miteinander zu kommunizieren und Stoffe auszutauschen, sind die Dictyosomen über Kanäle miteinander verbunden.

Die Golgi-Stapel sind asymmetrisch angeordnet. Man unterscheidet einen cis-, einen medialen-, einen trans-Golgi-Bereich und ein trans-Golgi-Netzwerk (TGN). Jedes dieser Subkompartimente besitzt zwar spezifische Marker, doch die Grenzen sind fließend.

Proteinmodifikation und -transport

Das glatte ER kommt vermehrt in Zellen vor, die sich auf die Lipidsynthese spezialisiert haben oder in steroidhormonproduzierenden Zellen. In den Membranen des ausgedehnten glatten ERs muss genügend Platz für die Enzyme der Cholesterinsynthese und der Hormonproduktion sein. In den Membranen sind Enzyme lokalisiert, die die Lipid-anteile der Lipoproteine synthetisieren. Lipoproteine sind notwendig, um Lipide über die Blutbahn zu anderen Zielzellen zu transportieren. Andere Enzyme sind an der Beseitigung fettlöslicher Medikamente und Giftstoffe beteiligt.

Wie funktionierten die Proteinherstellung und der Transport im ER?

Was mit dem Protein geschieht, hängt ab vom Ort seiner Synthese. Wir erinnern uns: Proteine, bzw. die Polypeptidkette, werden mittels Translation an den Ribosomen synthetisiert.

Enthält die wachsende Polypeptidkette ein ER-Signalpeptid, wird sie noch während der Synthese, d. h. co-translational, zusammen mit den Ribosomen zur Membran des ER dirigiert. Im Gegensatz dazu gelangen Proteine, die an freien Ribosomen im Zellplasma synthetisiert werden, erst nach vollendeter Synthese, post-translational, zum Zielkompartiment. Über diesen Transportvorgang ist bislang wenig bekannt. Es wird angenommen, dass sog. Hsp70-Proteine an diesem Prozess beteiligt sind.

Beim Transport eines Proteins mit ER-Signalpeptid zur ER-Membran sind vor allem zwei zelluläre Komponenten notwendig: ein Signalerkennungspartikel (SRP, signal recognition particle) und der entsprechende Rezeptor (Abb. 142).

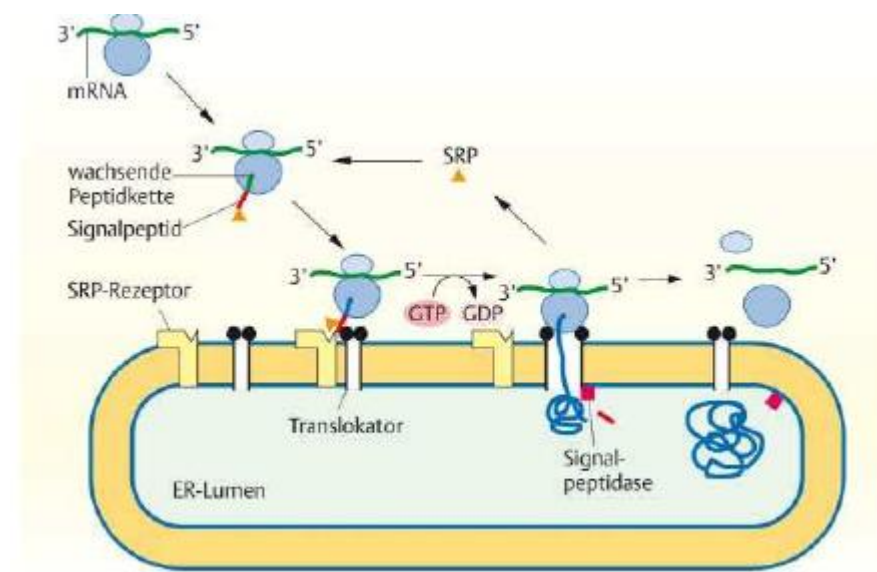


Abbildung 142 Transport von Proteinen an das raue ER

Zuerst bindet das SRP an das Signalpeptid sobald es am Ribosom entstanden ist. Die Proteinsynthese stoppt und das SRP dirigiert die Polypeptidkette zusammen mit dem Ribosom zur ER-Membran. Der Komplex aus SRP, mRNA, Ribosom und Polypeptidkette bindet an den SRP-Rezeptor, ein Membranprotein, das auf der Zellplasma-Seite der ER-Membran lokalisiert ist und Bindungsstellen für GDP/GTP und SRP besitzt. GTP ist mit dem ATP verwandt. Bindet das SRP an den SRP-Rezeptor, wird aus dem GTP GDP+P. Durch diese Spaltung des GTP wird das SRP wieder freigesetzt.

Der Transport der wachsenden Polypeptidkette durch die Membran hindurch erfolgt durch den Translokationsapparat. Bindet ein Ribosom an den Translokationsapparat, bildet dieser eine Pore, die sich nach erfolgter Proteinsynthese und Ablösen des Ribosoms schließt. Die zentrale Rolle im Translokationsapparat spielt der Sec61-Komplex: Er wird bei der Signalerkennung der Polypeptidkette benötigt, wenn diese in den Translokationsapparat eindringt, ist verantwortlich für die enge Bindung der Ribosomen und schafft mit seinen drei verschiedenen Proteinen die entsprechende Membrenumgebung für die Translokation der Polypeptidkette. Das Signalpeptid verbleibt im Translokationsapparat, während die wachsende Polypeptidkette kontinuierlich in Form einer Schleife durch die ER-Membran „hindurchsynthetisiert“ wird, bis sich das Protein vollständig im Lumen des ERs befindet, wo Chaperone die Faltung der Proteine unterstützen und überwachen. Chaperone sind Proteine, die neu

synthetisierte Proteine bei der Faltung unterstützen. Als letztes wird das Signalpeptid durch das Enzym Signalpeptidase abgespalten. Dazu ist eine zusätzliche Erkennungssequenz erforderlich, die dem Signalpeptid folgt.

Die Glykosylierung ist eine der wichtigsten Funktionen des ERs (Abb. 143).

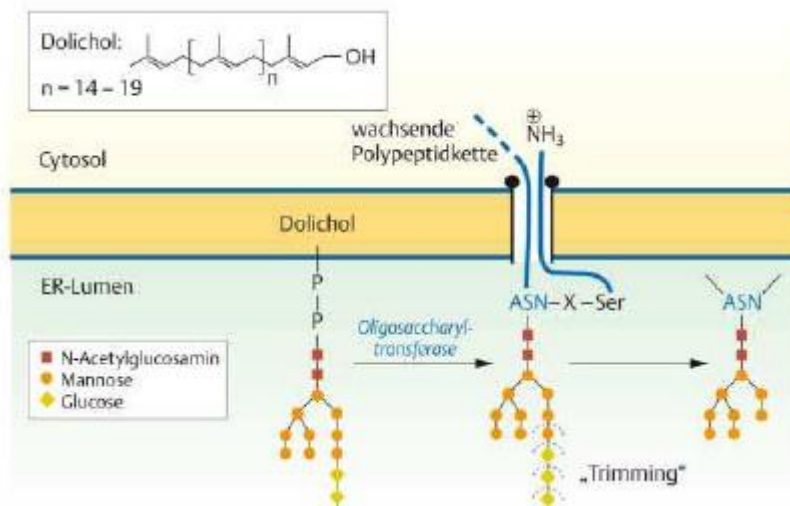


Abbildung 143 Glykolysierung von Proteinen

Die meisten Proteine, die das ER verlassen und weiter zu Golgi-Apparat, Lysosomen, Cytoplasmamembran oder Umgebung der Zelle transportiert werden, sind Glykoproteine. Glykoproteine sind eine Verbindung von Proteinen mit Kohlenhydraten. Die Verbindung mit Zuckermolekülen an Proteinen verläuft im ER mittels des Enzyms Oligosaccharyl-Transferase. Dieses Enzym überträgt einen kompletten Oligosaccharidkomplex auf die NH₂-Gruppe der Seitenkette eines Asparaginsäurerests. Hierfür gibt es spezielle Erkennungssequenzen aus drei Aminosäuren, an denen die Enzyme die Zuckermoleküle anheften können (Asp-X-Ser oder Asp-X-Thr, wobei X jede Aminosäure außer Prolin sein kann). Diese Übertragung wird auch als N-oder Asparagin-gekoppelte Glykosylierung bezeichnet und verläuft co-translational, also parallel zur Proteinsynthese. Durch weitere Modifikationen des Oligosaccharidkomplexes entsteht eine Vielfalt an Glykoproteinen. Bei den meisten Proteinen werden zunächst drei Glucose- und eine Mannosegruppe entfernt. Dieses sogenannte Trimming beginnt im ER und wird im Golgi-Apparat fortgesetzt.

Oligosaccharide üben am Protein verschiedene Funktionen aus: Sie können das Protein vor Abbau schützen, das Protein im ER festhalten bis es korrekt gefaltet ist oder als Transportsignal dafür sorgen, dass das Protein in Vesikel verpackt und zum entsprechenden Kompartiment transportiert wird. Erscheinen Oligosaccharide auf der Zelloberfläche, so sind sie Teil der Glykokalyx und können an der Zell-Zell-Erkennung beteiligt sein.

Die Lipidsynthese findet an den Membranen des glatten ERs von Tier- und Pilzzellen statt. Sie betrifft Lipide, die zum Membranaufbau benötigt werden (Phospholipide, Cholesterin), Reservelipide (Triacylglycerol) und andere lipophile Verbindungen wie Steroide.

Die Phospholipidsynthese findet fast ausschließlich an der Zellplasma-Seite des ERs statt, da alle notwendigen Enzyme in der ER-Membran lokalisiert sind und ihre aktiven Zentren zum Zellplasma hin ausgerichtet haben. Damit sich die gesamte Lipiddoppelmembran vergrößern kann und das Wachstum nicht auf ein Membranblatt beschränkt bleibt, sind Phospholipid-Translokasen („Flippasen“) notwendig, die Phospholipidmoleküle auch zur luminalen Seite des ERs bringen. Diese Flippasen sind spezifisch für die entsprechende Kopfgruppe des Phospholipids.

Zwischen ER und Cytoplasmamembran, Golgi-Apparat, Lysosomen und Endosomen findet der Transport von Proteinen und Lipiden über Transportvesikel statt (Sekretions- und Endocytoseweg) statt.

Bedingung für den Austritt von Proteinen aus dem ER ist Vollständigkeit und korrekte Faltung. Andernfalls wird das Protein ausgesondert und abgebaut. Dabei spielen Chaperone eine wichtige Rolle, die an unvollständig oder falsch gefaltete Proteine binden und deren Weitertransport verhindern.

Proteine, die nicht weiter transportiert werden, sondern im ER-Lumen verbleiben, werden als ER-ständige Proteine bezeichnet. ER-ständige Proteine sorgen z. B. für die richtige Faltung vieler Proteine, die in das Lumen des ERs gelangen. Ein Beispiel für ein solches ER-ständiges Protein ist die Protein-Disulfid-Isomerase (PDI), die die Bildung von Disulfid-(S-S)-Brücken katalysiert. Ein weiteres ER-ständiges Protein ist das Bindeprotein (BiP), das mit den Hsp70-Proteinen, also Hitzeschockproteinen, verwandt ist und als Chaperon dient. Durch die Bindung an BiP kann die Ausschleusung ER-ständiger Proteine aus dem ER sowie deren Aggregation verhindert werden.

Funktionen des Golgi-Apparates

Der Golgi-Apparat empfängt Proteine und Lipide aus dem ER. Die posttranslationale Proteinmodifizierung, die im ER begonnen hat, wird hier fortgeführt.

Die Übertragung von Oligosacchariden auf bestimmte Asparaginsäurereste der Proteine sowie deren Zurechtschneiden findet schon im ER statt. Nach dieser N-Glykosylierung können sich weitere Modifizierungen im Golgi-Apparat anschließen.

Im Golgi-Apparat findet neben der Glykosylierung von Proteinen auch eine Glykosylierung von Lipiden statt. Dabei entstehen Glykolipide wie Cerebroside und Ganglioside.

Nachdem die Proteinmodifikationen im Golgi-Apparat beendet sind, schließt sich im Trans-Golgi-Netzwerk das Sortieren der Proteine an. Der Transport zu den anderen Kompartimenten des Endomembransystems erfolgt durch kontinuierliche Abschnürung und Fusion von Transportvesikeln. Diese Vesikel verbinden die Kompartimente miteinander, indem sie sich vom Donorkompartiment abschnüren, ihren Inhalt in Richtung Akzeptorkompartiment transportieren, dort mit der Membran verschmelzen und ihren Inhalt freisetzen.

Bei der Vesikelknospung bildet sich eine Hülle aus Proteinen um den knospenden Membranbereich. Es gibt drei verschiedene, gut charakterisierte Hüllentypen, COPI, COPII und Clathrin, sowie zahlreiche andere, über die noch sehr wenig bekannt ist. Vesikel, die COPI und COPII als Hülle tragen, sind am Transport zwischen ER und Golgi-Apparat sowie innerhalb des Golgi-Apparats beteiligt. Dabei transportieren COPII-Vesikel Material vom ER zum Golgi, während COPI-Vesikel den Transport innerhalb des Golgi-Apparates und vom Golgi zum ER erledigen (Abb. 144).

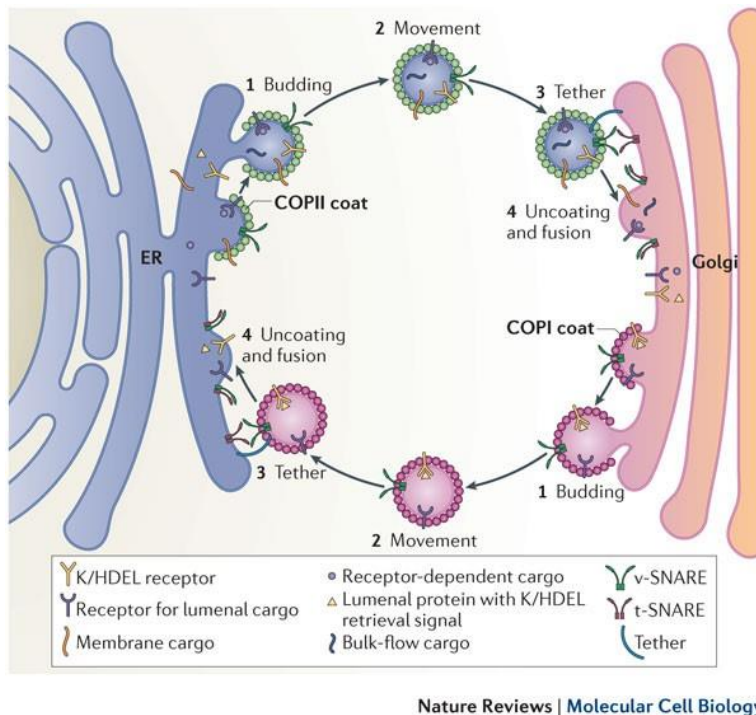


Abbildung 144 Transport im Golgi-Apparat

Der Knospungsvorgang wird an Vesikeln durch die Zusammenlagerung der COP-Hüllproteine und Arf, einem kleinen G-Protein, initiiert. Im Cytoplasma liegen diese Komponenten inaktiv vor. Arf trägt gebundenes GDP und auf ein Signal hin wird Arf durch Austausch von GDP gegen GTP aktiviert und lagert sich an der Donormembran an. An dieser Stelle lagern sich die Hüllproteine zusammen und durch die Bildung einer dichten Proteinhülle wird die Knospung vorangetrieben bis ein COP-umhüllter Transportvesikel entsteht. Gelangt das Vesikel an die Zielmembran, muss zunächst die Hülle entfernt werden. Das Signal dafür ist die GTP-Hydrolyse an Arf, worauf die Desintegration der gesamten Proteinhülle folgt.

Wie erkennen Vesikel ihr Ziel? Die Spezifität des gerichteten vesikulären Transportes, Proteintargeting, setzt voraus, dass alle Arten von Transportvesikeln in der Zelle auf ihrer Oberfläche Marker besitzen, die Auskunft über ihre Beladung und ihren Ursprung geben. Diese Marker wiederum müssen von komplementären Rezeptoren auf der entsprechenden Zielmembran erkannt werden. An diesem Erkennungsmechanismus ist eine Familie von Membranproteinen beteiligt, die SNAREs (soluble NSF attachment proteins). In den Vesikeln gibt es die v-SNAREs und auf der Zielmembran die t-SNAREs. Treffen beide aufeinander (es bildet sich ein trans-SNARE-Komplex), kommt es zu einer Konformationsänderung und nähert die beiden Membranen so sehr aneinander an, dass es zur Fusion kommt. Der gerichtete vesikuläre Transport wird zusätzlich durch rab-Proteine kontrolliert. Sie prüfen, ob jedes Vesikel mit der richtigen

Ziellmembran fusioniert. Diese kleinen G-Proteine sind nicht nur am vesikulären Transport, sondern auch an Endo- und Exocytosevorgängen beteiligt. Nach erfolgter Fusion von Vesikel und Ziellmembran befindet sich der trans-SNARE-Komplex vollständig in der Ziellmembran. Hier interagiert er mit einem Fusionsprotein namens NSF (N-Ethylmaleimid-sensitiven Fusionsprotein) und weiteren Fusionsproteinen. Die Interaktion von NSF mit dem trans-SNARE-Komplex führt unter ATP-Verbrauch zur Dissoziation der trans-SNARE-Komplexe, sodass die v-SNAREs wieder frei vorliegen. Sie werden aus der Ziel-membran zurückgeführt und stehen dann für neue Fusionsrunden zur Verfügung (Abb. 145).

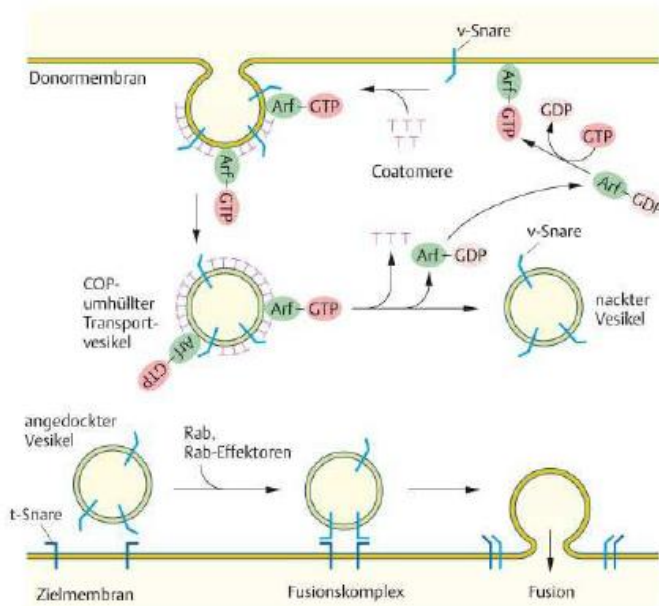


Abbildung 145 Vesikelbildung im GA

Rolle des Cytoskeletts beim Transport

Das Cytoskelett spielt beim Transport von Vesikeln eine wichtige Rolle, allen voran die Mikrotubuli, über deren Aufbau und Rolle wir bei der Mitose einiges erfahren haben.

Sie sind dabei in der Lage, Vesikel durch die Zelle zu transportieren. Die Vesikel enthalten meist verschiedenste Signalstoffe, die zu den unterschiedlichen Zellorganellen transportiert werden.

Die Mikrotubuli legen eine Art „Schienensystem“ für die Vesikel, durch das sie sich bewegen können. Die Vesikel binden sich dabei über sogenannte Motorproteine (Kinesin, Dynein, Myosin) an die Mikrotubuli. Ein solches Motorprotein besteht grundsätzlich aus einer sogenannten Kopfregion mit zwei Köpfen und einer Schwanzregion. Die Kopfregion ist in der Lage, an die Mikrotubuli zu binden und Energie in Form von ATP zu produzieren. Die Schwanzregion bindet die Vesikel an sich. Durch das Spalten von ATP zu ADP wird Energie frei, welche die räumliche Anordnung der Kopfregion verändern kann. Das Motorprotein bewegt sich so

schrittweise nach vorne, wobei immer einer der beiden Köpfe an den Mikrotubulus gebunden ist (Abb. 146).

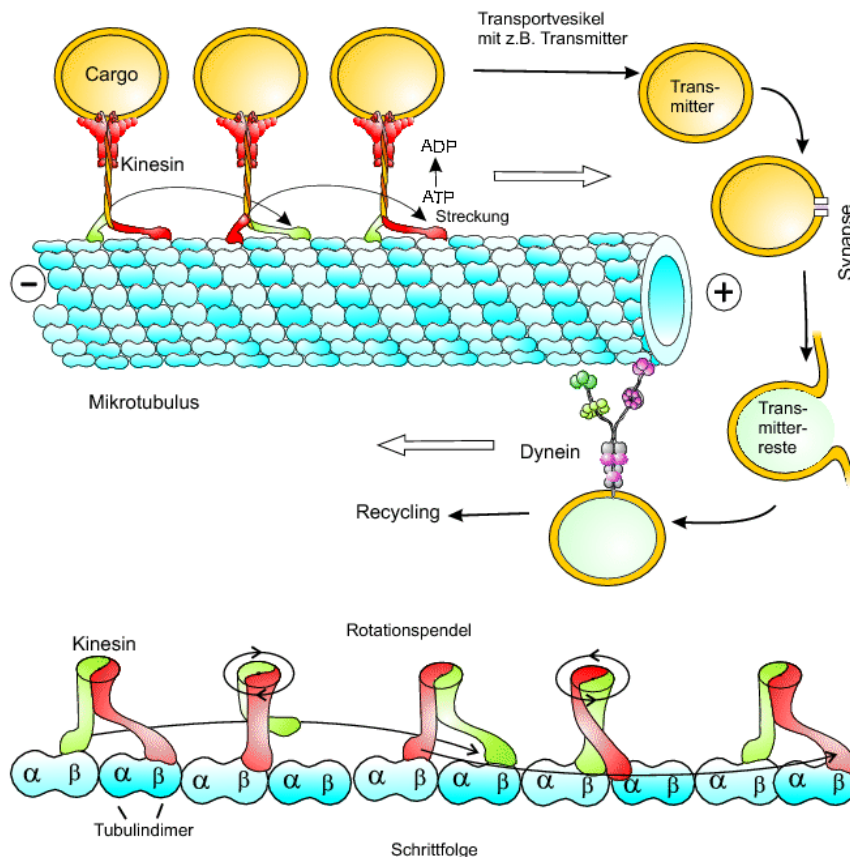


Abbildung 146 Vesikeltransport

Einen wichtigen Vorgang, den man sich merken sollte, ist die sogenannte Exocytose. Bei ihr „verschmelzt“ zuerst das Vesikel mit der Zellmembran und gibt danach ihren Inhalt frei.

Die Exozytose (Abb. 147) dient dazu, Abfall- und Nebenprodukte von Stoffwechselvorgängen aus der Zelle zu schleusen. Außerdem können die Zellen über Exozytose auch nützliche Stoffe nach außen abgeben. Das können zum Beispiel Neurotransmitter oder Hormone sein, die eine reibungslose Kommunikation zwischen mehreren Zellen ermöglichen.

Prinzipiell kann man zwischen einer konstitutiven und einer stimulierten Exozytose unterscheiden. Bei der stimulierten Exozytose ist ein Reiz (z.B. Hormon) für die Ausschüttung der Substanzen verantwortlich, was bei der konstitutiven Exozytose nicht der Fall ist. Sie erfolgt ohne äußere Aktivierung und sorgt zum Beispiel dafür, dass die Biomembranen erweitert werden können.

Die Exozytose steht im Gegensatz zu der Endozytose, die eine Aufnahme von Vesikeln in das Zellinnere beschreibt.

Die Aufnahme von großen Partikeln oder ganzen Zellen nennt man Phagocytose. Sie ist bezeichnend für Zellen mit hoher phagocytotischer Aktivität (z.B. Amöben, viele Ciliaten, Makrophagen, neutrophile Granulocyten).

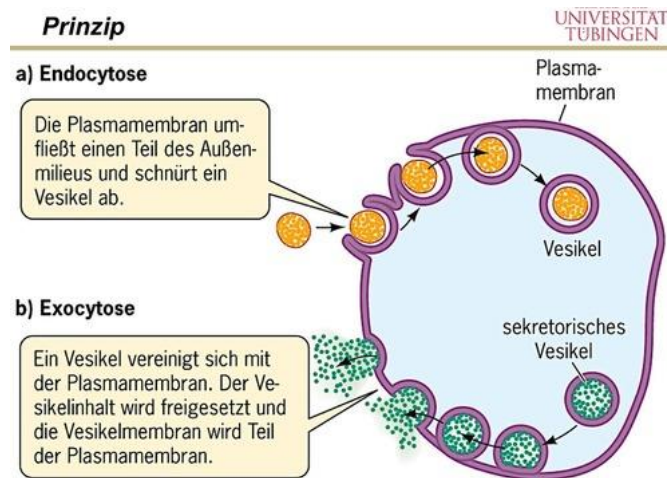


Abbildung 147 Endo- und Exocytose

Lysosomen

Die dritte Funktion des Golgi-Apparates ist die Bildung von sogenannten Lysosomen. Diese sind Zellorganellen in eukaryotischen Zellen. Grundsätzlich ähneln sie in ihrem Aufbau den Vesikeln sehr. Sie enthalten verschiedenste Enzyme, die die Verdauung unterstützen können.

In den pflanzlichen Zellen hat der Golgi-Apparat noch eine weitere wichtige Funktion. Hier produziert er verschiedenste Polysaccharide, also Kohlenhydrate. Diese kann die Pflanzenzelle verwenden, um ihre Zellwand aufzubauen.

Ein Lysosom hat eine rundliche Form und ist von einer Membran umgeben.

Der pH-Wert in seinem Inneren (= Lumen) ist sauer und liegt in etwa in einem Bereich um 4,5 – 5. Das Cytosol, in dem sich die Lysosomen bewegen, hat einen neutralen pH-Wert von ca. 7,2 (Abb. 148).

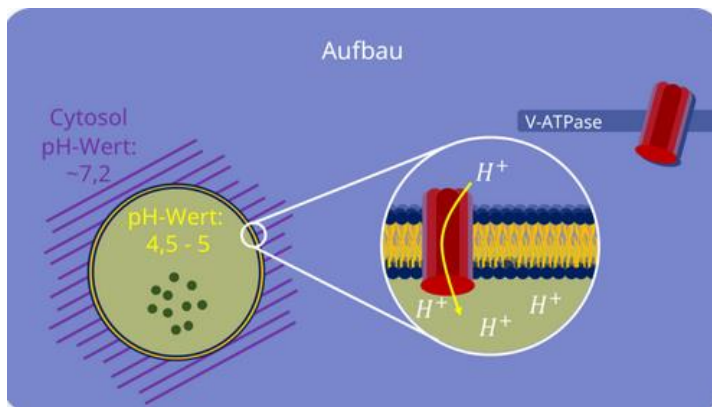


Abbildung 148 Lysosom

In diesem sauren Milieu weisen die verschiedenen Enzyme wie zum Beispiel Lipasen, Proteasen und Nukleasen eine hohe Aktivität auf. Diese sogenannten hydrolytischen Enzyme sind in der Lage, Lipide, Proteine, Nukleinsäuren und Polysaccharide durch eine Hydrolyse zu spalten. Ohne die saure Umgebung könnten diese Enzyme nicht effektiv arbeiten, die genannten Stoffe nicht spalten und somit könnte das Lysosom die Stoffe nicht verdauen.

Um den niedrigen pH-Wert zu gewährleisten, benötigt das Lysosom das Transportprotein V-Typ-ATPase, welches in die Lysosommembran eingebettet ist. V-Typ-ATPase ist in der Lage, ATP zu hydrolysieren und dadurch Protonen (H^+) vom Cytosol durch die Membran zu „pumpen“. Da der pH-Wert mit steigender Anzahl an H^+ -Ionen niedriger wird, entsteht so eine durchgängig saure Umgebung im Lysosom (Abb. 149).

Für die Entstehung der Lysosomen ist im Voraus eine Bildung der lysosomalen Enzyme notwendig. Diese Bildung findet im rauen Teil des Endoplasmatischen Retikulums statt. Dort werden sie außerdem in Vesikel gebunden und zum Golgi-Apparat transportiert.

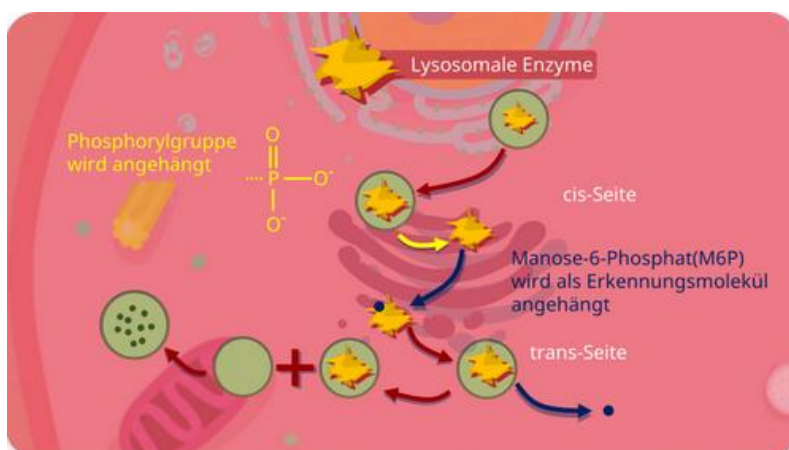


Abbildung 149 Lysosom-Entstehung

Auf der cis-Seite des Golgi-Apparates findet eine Vorbereitung für den Weitertransport der lysosomalen Enzyme statt. Dabei wird eine Phosphorylgruppe (PO_3^{2-}) an das Enzym angehängt (= Phosphorylierung). Danach wird diesen phosphorylierten

Enzymen auf der trans-Seite des Golgi-Apparates ein Marker namens Mannose-6-Phosphat (M6P) gegeben. Dieser Marker dient als Erkennungsmolekül für lysosomale Enzyme (Abb. 149).

Im nächsten Schritt werden die lysosomalen Enzyme in Vesikel verpackt. So entstehen die sogenannten primären Lysosomen. Der M6P Marker löst sich ab und kann wiederverwendet werden.

Alle Proteine, die im Golgi-Apparat nicht phosphoryliert wurden, erhalten den M6P-Marker auch nicht. So erkennt die Zelle, dass es sich dabei nicht um lysosomale Proteine handelt. Diese Proteine werden dann in Transportvesikel verpackt und durch Exocytose aus der Zelle herausbefördert.

Die wichtigste Funktion eines Lysosoms besteht darin, zelleigene und zellfremde Stoffe abzubauen. Dafür schließt es den zu verdauenden Stoff (Protein, ...) in seinem Innenraum ein. Ein Lysosom ist auch in der Lage, zelleigene Stoffe, wie zum Beispiel Teile von Organellen oder des Cytosols, zu verdauen (= Autophagie). Diese können danach bei Bedarf sogar wiederverwertet werden. Das Lysosom ist also hier eine Art Reparaturvesikel für unsere Zellen.

Peroxisomen

Als letztes seien die Peroxisomen erwähnt (Abb. 150). Ein Peroxisom (eng. peroxisome) ist ein Zellorganell in eukaryotischen Zellen. Peroxisomen verdanken ihren Namen ihrer Fähigkeit, Wasserstoffperoxid abzubauen.

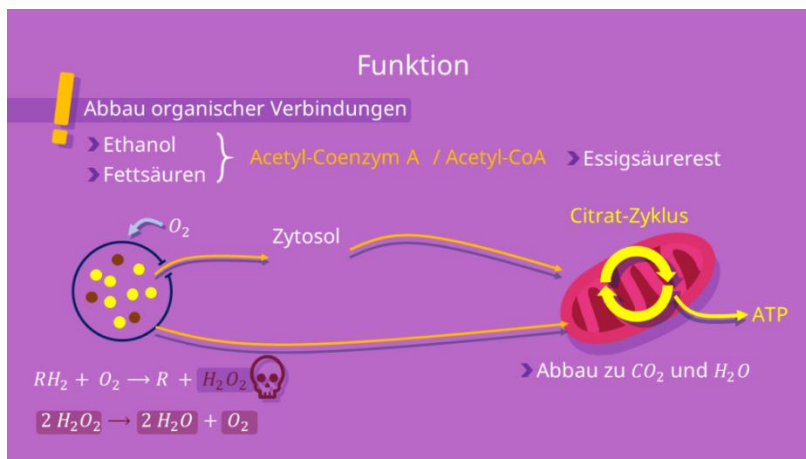


Abbildung 150 Peroxisom und seine Funktionen

Eine Funktion des Peroxisoms ist der oxidative Abbau organischer Verbindungen. Das können zum Beispiel Ethanol oder Fettsäuren sein. Insbesondere Ethanol wird vermehrt in den Leberzellen durch Peroxisomen abgebaut.

In tierischen Zellen werden die Fettsäuren und das Ethanol zum sogenannten Acetyl-Coenzym A, kurz CoA, oxidiert. Diesen haben wir bei der Zellatmung kennengelernt.

Jedoch benötigen die Enzyme des Peroxisoms für den Abbau dieser organischen Verbindungen Sauerstoff als zusätzliches Substrat. Durch diesen Sauerstoff kann sich in der Zelle Wasserstoffperoxid (H_2O_2) als Nebenprodukt bilden.

Wasserstoffperoxid gilt als Zellgift für alle Zellorganellen und insbesondere für das Cytoplasma.

Dieses Zellgift kann das Peroxisom durch seine nächste Funktion auflösen. Die Katalase in seinem Lumen ist in der Lage, Wasserstoffperoxid in die ungefährlichen Moleküle Wasser (H_2O) und Sauerstoff (O_2) aufzutrennen.

Somit dient das Peroxisom als Reaktionsraum, um die Zelle zu entgiften.

Die dritte Funktion der Peroxisomen bezieht sich auf die Synthese der Myelinscheide von Nervenzellen. Myelin ist eine fetthaltige Biomembran. Wenn diese die Axone der Nervenzellen von Wirbeltieren umwickelt, entsteht die Myelinscheide.

Kapitel 16: Zellgewebe und Zellverknüpfung

Lebewesen können wie gesagt als Einzeller oder Mehrzeller vorkommen. Wenn viele Zellen zusammen einen Verbund eingehen, spricht man von Gewebe. Die meisten Tiere, Pilze und Pflanzen sind Vielzeller, bestehen also aus einem Verband mehrerer Zellen.

Gewebetypen bei Tierzellen

Ein Zellgewebe ist eine Ansammlung differenzierter Zellen einschließlich ihrer extrazellulären Matrix.

Die Zellen eines Gewebes besitzen ähnliche Funktionen und erfüllen gemeinsam die Aufgaben des Gewebes. Alle Anteile der meisten Vielzeller lassen sich einem Gewebe zuordnen, bzw. sie sind von einem Gewebetyp produziert worden. Hier soll erstmal eine kurze Auflistung erfolgen:

Die tierischen Gewebe (Abb. 151) sind:

1. a) Epithelgewebe: Zellschichten, die alle inneren und äußeren Oberflächen bedecken. Es wird grob in Oberflächen- und Drüsenepithelien gegliedert.
2. b) Binde- und Stützgewebe: Gewebe, das für strukturellen Zusammenhalt sorgt und Zwischenräume füllt (hierzu gehören auch Knochen, Knorpel und Fettgewebe) und im weitesten Sinne weitere spezialisierte Gewebe (z. B. das Blut) hervorbringt.
3. c) Muskelgewebe: Zellen, die durch kontraktile Filamente für aktive Bewegung spezialisiert sind.
4. d) Nervengewebe: Zellen, aus denen Gehirn, Rückenmark und periphere Nerven aufgebaut sind.

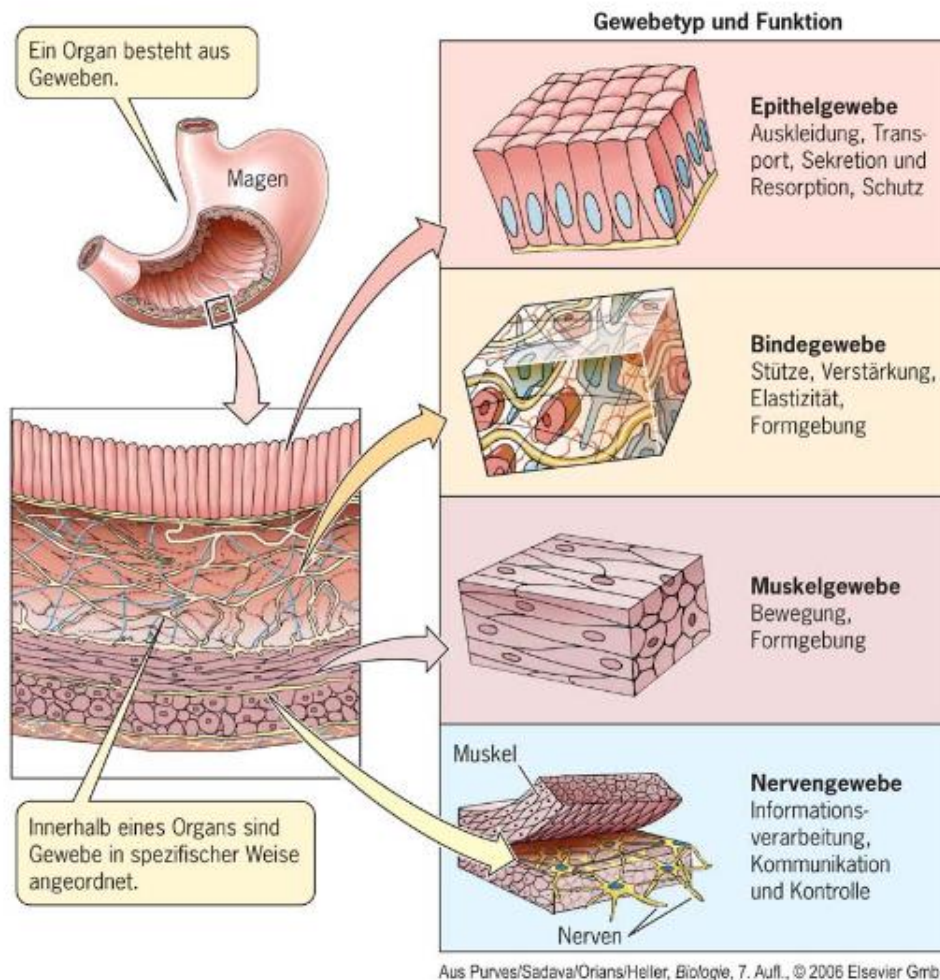


Abbildung 151 tierische Gewebetypen

Oberflächenstrukturen und extrazelluläre Matrix

Die Oberfläche einer Zelle erfüllt verschiedene Aufgaben. Sie tritt in Kontakt mit anderen Zellen, grenzt sie aber gleichzeitig von anderen ab, regelt den Stofftransport in die Zelle und aus der Zelle hinaus und hat Rezeptoren für Signalmoleküle. Aufgrund ihrer verschiedenen Aufgaben sind Zelloberflächen außerordentlich vielgestaltig. Zu den Oberflächendifferenzierungen zählen so verschiedene Strukturen wie die Axone und Dendriten von Nervenzellen, Mikrovilli, Flagellen, Cilien und sonstige Zellfortsätze. Direkt mit der Cytoplasmamembran assoziiert ist die Glykokalyx. Darüber hinaus umgeben sich die Zellen mit extrazellulärem Material. Dieses wird von ihnen selbst sezerniert und bei den Tieren als extrazelluläre Matrix (Abb. 152), bei Pflanzen, Pilzen und Prokaryoten als Zellwand bezeichnet.

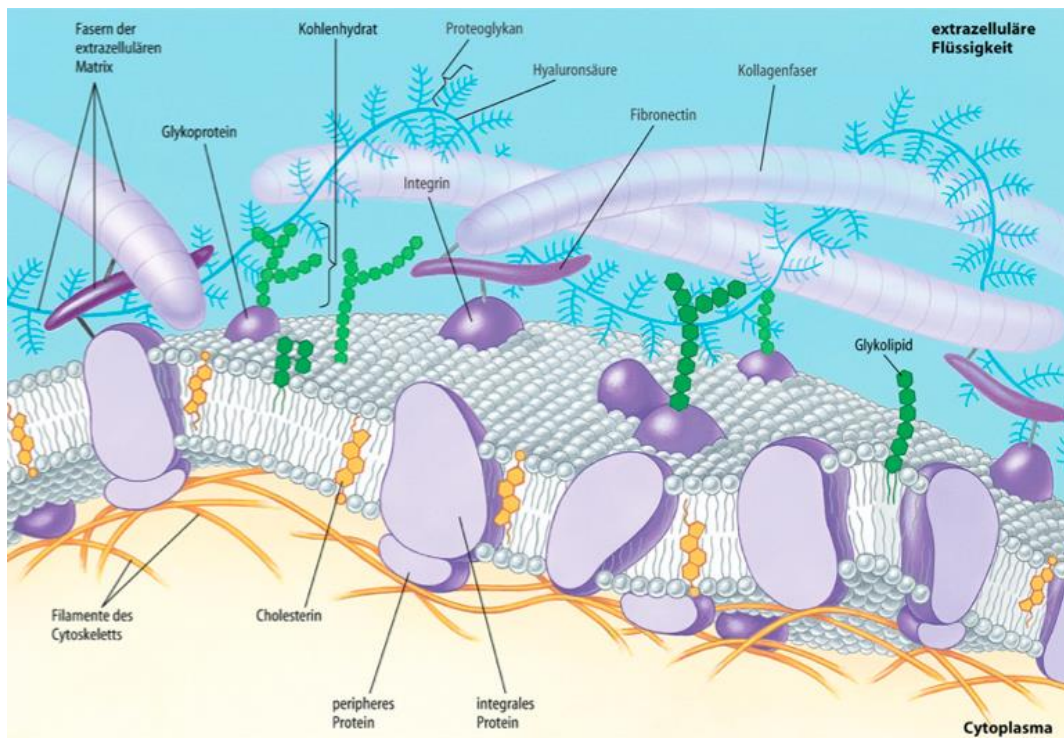


Abbildung 152 extrazelluläre Matrix

Mit der Cytoplasmamembran sind eine Reihe von Glykoproteinen und Glykolipiden direkt oder indirekt assoziiert, deren teilweise sehr komplexen Oligosaccharid-Seitenketten ausschließlich auf der Außenseite exponiert sind. Der Besatz mit Glykoproteinen und Glykolipiden kann sehr dicht sein und einen regelrechten Zuckermantel um die Zellen bilden. Er wird als Glykokalyx bezeichnet. Der Zuckeranteil kann bis zu 10 % der Masse von Cytoplasmamembranen ausmachen. Die an der Oberfläche exponierten Strukturen sind für Zell-Zell-Erkennung, Zellkontakte und Signalwahrnehmung verantwortlich. Der Zellerkennung dienen beispielsweise Glykolipide, wie die Blutgruppenantigene des AB0-Systems oder die MHC-Komplexe, welche für die Gewebeverträglichkeit bei Organtransplantationen von entscheidender Bedeutung sind.

Die extrazelluläre Matrix (EZM) ist der Gewebeanteil, der zwischen den Zellen im sogenannten Interzellularraum bei tierischen Zellen liegt. Die extrazelluläre Matrix tierischer Vielzeller setzt sich aus diversen Komponenten zusammen, die in zwei große Gruppen eingeteilt werden: Grundsubstanz und Fasern. Die Fasern dienen der Formgebung und bestehen aus Kollagen oder Elastin, die in die Grundsubstanz eingebettet sind. Die Grundsubstanz besteht aus verschiedenen sog. Proteoglykankomplexen, Glucosaminoglykanen (wie Hyaluron) und/oder Glykoproteinen. Das Verhältnis der Grundsubstanz zum Faseranteil schwankt je nach Lokalisation ebenso wie der Anteil der extrazellulären Matrix am Gewebe insgesamt, bedingt durch dessen jeweilige Funktion.

Weitere wichtige Bestandteile der EZM sind Verbindungsmoleküle, die verschiedene Komponenten der EZM untereinander und diese mit der Zelloberfläche verbinden. Zwei der Wichtigsten sind Fibronectin und Laminin. Fibronectin enthält Bindungsstellen für Kollagen und Heparin. Durch alternatives Spleißen der mRNA entstehen verschiedene Fibronectin-Formen, die bei einer Reihe von biologischen

Prozessen eine Rolle spielen. Dazu gehört die Blutgerinnung, Wundheilung oder die Wanderung der Neuralleistenzellen während der Embryonalentwicklung.

Laminin ist ein charakteristischer Bestandteil der Basallaminae und ähnlich wie Fibronectin enthält auch dieses Protein Bindungsstellen für verschiedene Moleküle wie Kollagen IV und Heparansulfat.

Die wichtigsten Rezeptoren an der Zelloberfläche für Komponenten der EZM sind die Integrine, die an verschiedene Liganden binden.

Zwischen den Zellen tierischer Gewebe gibt es eine Reihe von Zell-Zell-Verbindungen: Gap Junctions, Tight Junctions und Desmosomen.

Gap Junctions

Gap Junctions (Abb. 153) kommen vor allem in Muskel- und Nervengeweben tierischer Organismen vor, da dort eine schnelle Kommunikation erforderlich ist.

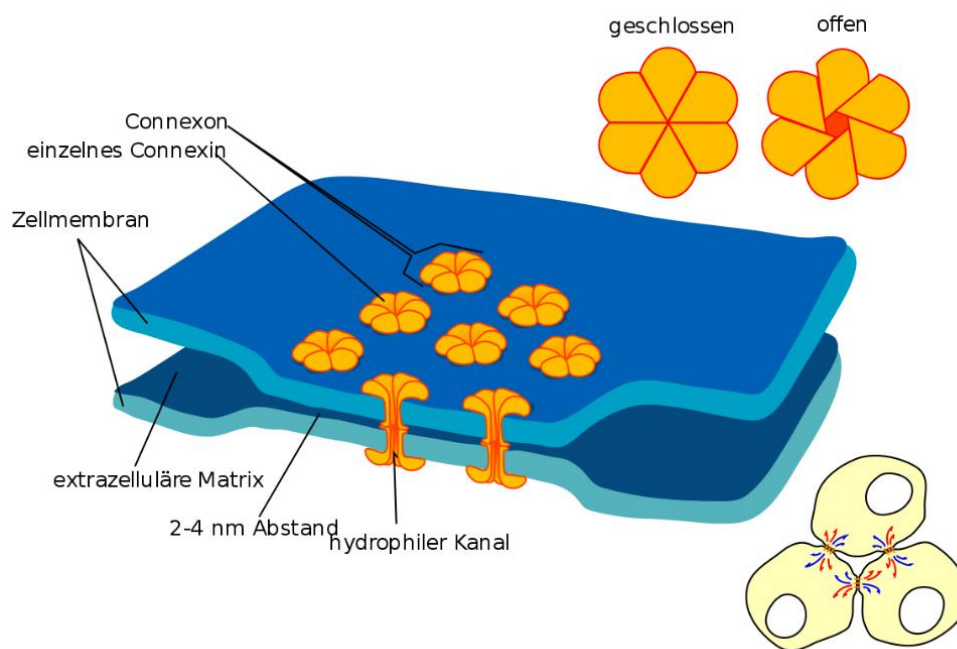


Abbildung 153 Gap Junction

Es sind Proteinkanäle zwischen den Plasmamembranen zweier angrenzender tierischer Zellen. Sie ermöglichen somit eine direkte Verbindung des Cytoplasmas der beteiligten Zellen. Sie dienen zur Kommunikation zwischen Zellen, indem chemische Substanzen oder elektrische Signale weitergeleitet werden. sie werden etwa 2,7 nm groß.

Sie durchspannen die jeweiligen Plasmamembranen der benachbarten Zellen und den dazwischen liegenden Interzellularraum.

Das beteiligte Kanalprotein heißt Connexin. Sechs Connexine bilden nun jeweils eine ringförmige Struktur – das Connexon. Lagern sich jetzt jeweils zwei Connexone der benachbarten Plasmamembranen aneinander, bildet sich eine durchgehende Pore. Diese Pore entspricht einer Gap Junction.

Durch eine unterschiedliche Zusammensetzung der Connexine kann die Durchlässigkeit oder Leitfähigkeit der Kanäle jeweils variieren. Bei uns Menschen gibt es beispielsweise über 20 verschiedene Connexin-Gene. Der Durchmesser der Poren beträgt circa 1,5 nm, weshalb großen Molekülen wie Proteinen der Zutritt verwehrt bleibt.

Im Vergleich zu den anderen Kanalsystemen in der Zelle unterscheiden sich die Gap Junctions vor allem darin, dass hier zwei Membranen anstatt einer durchzogen werden.

Die Hauptfunktion der Gap Junctions liegt in der Kommunikation zwischen benachbarten Zellen.

Eine Möglichkeit über Gap Junctions zu kommunizieren bietet die elektrische Kommunikation. Darunter versteht man eine schnelle Weiterleitung von elektrischen Signalen in Form von Ionenströmen. In Nervenzellen in der Netzhaut (Retina) am Auge oder im Herzen dienen sie als sogenannte elektrische Synapsen. Diese ermöglichen eine schnelle und gleichmäßige Ausbreitung von elektrischen Potentialen.

Neben der elektrischen Kommunikation kann auch eine chemische Kommunikation über Gap Junctions stattfinden. Signalmoleküle, sogenannte second messenger wie Calcium-Ionen (Ca^{2+}), können die Gap Junctions durchqueren. Dadurch wird sichergestellt, dass diese Signalmoleküle in benachbarten Zellen in gleicher Konzentration vorhanden sind. Das ermöglicht eine koordinierte Reaktion der Zellen.

Außerdem können kleine Moleküle wie ATP, Glucose oder Aminosäuren durch die Kommunikationskanäle passieren. Sie sind für den Stoffwechsel in unseren Zellen sehr wichtig.

Gap Junctions sind nur in vielzelligen Tieren zu finden. Sie kommen vor allem dort vor, wo eine schnelle Kommunikation erforderlich ist, wie bei Muskel- und Nervenzellen. Zusätzlich unterstützen sie in Drüsen wie Leber und Bauchspeicheldrüse die Ausscheidung von Verdauungssekreten.

Desmosomen

Die Desmosomen sind in tierischen Zellen vertreten (Abb. 154). Man findet sie insbesondere in den Zellen des Herzmuskels und den Epithelzellen.

Desmosome

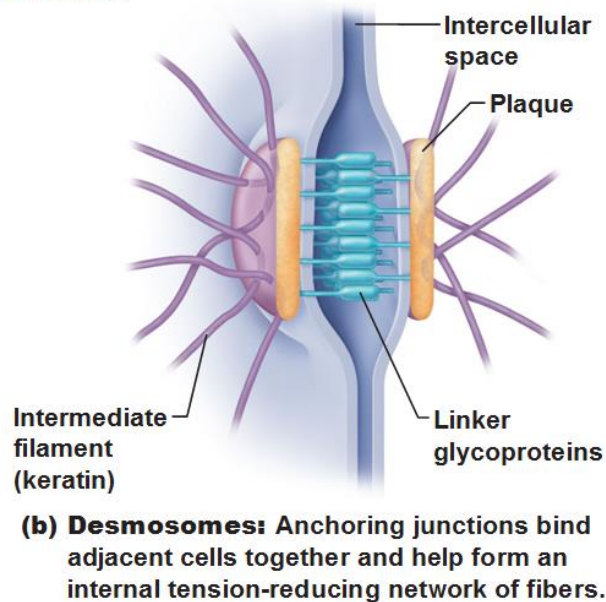


Abbildung 154 Desmosom

Ein Desmosom ist eine kugelförmige Struktur in der Zellmembran, vergleichbar mit Druckknöpfen.

Die wichtigste Funktion der Desmosomen ist die Herstellung einer Verbindung zwischen zwei Zellen, um dadurch den mechanischen Zusammenhalt zu erhöhen.

Die Zellstrukturen, welche eine Zelle nicht mit einer anderen Zelle, sondern mit dem Extrazellulärraum verbinden, werden als Hemidesmosomen bezeichnet.

Ein Desmosom setzt sich aus einem Netzwerk von Proteinen zusammen. All diese Proteine sind an der sogenannten Plaque befestigt. Man kann sich die Plaque als eine Art Verbindungsstelle im Desmosom vorstellen, die vor allem aus Ankerproteinen wie zum Beispiel Plaktoglobin und Plaktophilin besteht. Wie der Name schon sagt, dienen Ankerproteine als Verbindungsstücke für weitere Proteine.

An die Ankerproteine der Seite der Plaque, die zum Extrazellulärraum zeigt, sind die Cadherine wie zum Beispiel Desmoglein und Desmocollin befestigt. Diese länglichen Glykoproteine gehören zu den Adhäsionsproteinen. Ihre Aufgabe ist die Stabilisierung der Zell-Zell-Kontakte, indem sie sich mit den Cadherinen der anderen Zellen verbinden.

Die Seite der Plaque, die zum Zellinnenraum zeigt, enthält vor allem das Desmoplakin. Dieses spezielle Protein kommt nur in Desmosomen vor. Es bindet sich an die Ankerproteine und stellt eine Verbindung zu den Intermediärfilamenten des

Cytoskeletts her. Diese Verbindung wird insbesondere in den Epithelzellen über sogenannte Fibrillenbündel aus den Proteinen Keratin oder Cytokeratin (Zytokeratin) erzeugt. Diese Fibrillenbündel aus Keratin werden auch als Tonofilamente bezeichnet.

Durch die Verknüpfung mit dem Cytoskelett können sich die Desmosomen an ihrem Ort an der Zelloberfläche stabilisieren.

Die Hauptaufgabe der Desmosomen ist die Verankerung der Zellen in Tieren untereinander (Zell-Zell-Kontakte).

Die Desmosomen stellen eine sehr enge Verbindung zwischen den Zellen her. Der Abstand der Plasmamembranen im Desmosom ist nur ungefähr 15-25 nm groß. Dadurch ist der gesamte Zellverbund gegenüber äußeren mechanischen Einwirkungen (Zug-/Scherkräfte) relativ stabil. Aus diesem Grund befinden sie sich insbesondere in Zellen mit hoher mechanischer Belastung (z.B. Herzmuskel).

Doch nicht nur die Stabilität des Zellverbundes gegen äußere Einflüsse kann durch Desmosomen erhöht werden. Indem sie eine Verbindung zwischen der Zellmembran und dem Cytoskelett beispielsweise über die Tonofilamente herstellen, können sie auch eine einzelne Zelle relativ gut in ihrer Form halten.

Die Hemidesmosomen sind eng mit den Desmosomen verwandt. Wie auch ein Desmosom, kommt ein Hemidesmosom sehr oft in Epithelzellen vor.

Wenn man ihren Gesamtaufbau in einem Elektronenmikroskop betrachtet, wird auffallen, dass Hemidesmosomen wie halbe Desmosomen aussehen. Sie enthalten ebenfalls eine Plaque, an die nun aber keine Desmogleine, sondern Kontaktproteine namens Integrine gebunden sind. Außerdem sind sie gleichermaßen intrazellulär mit den Intermediärfilamenten des Cytoskeletts über Keratinfilamente verknüpft.

Ihre Funktion besteht nun jedoch nicht mehr darin, zwei Zellmembranen ähnlicher Zellen zu verbinden. Hemidesmosomen sind in der Lage, die Zellmembran der Epithelzellen mit den extrazellulären Basallamina zu verbinden. Das ist die Proteinschicht, die an den basalen Teil der Epithelzelle anschließt. Diese Basallamina können sich über sogenannte Laminin Proteine, mit den Integrinen der Hemidesmosomen verbinden.

Tight Junctions

Tight Junctions kommen in tierischen Zellen vor allem im Darm, Magen und in der Blut-Hirn-Schranke vor. Sie sorgen für eine Verbindung zwischen Epithelzellen (Abb. 155).

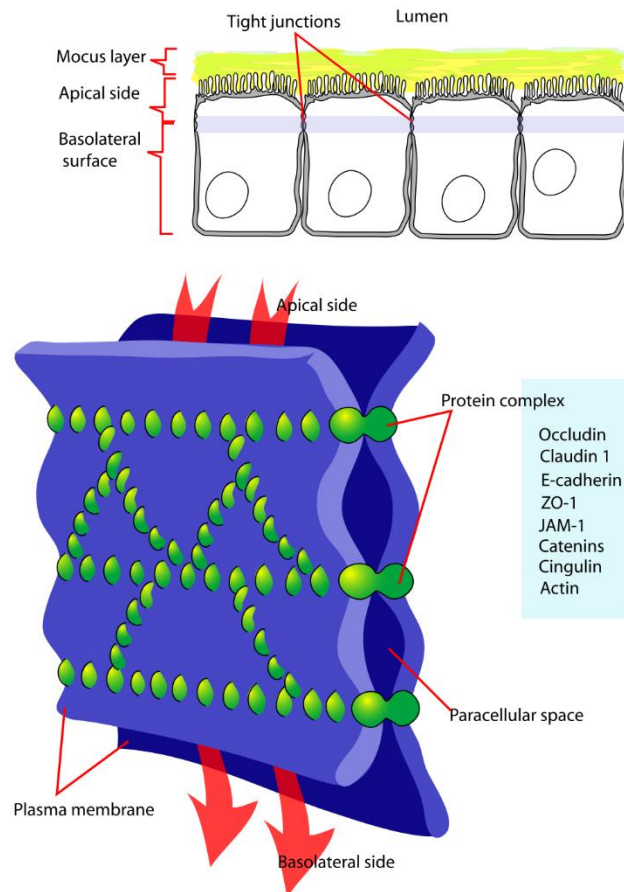


Abbildung 155 Tight Junction

Tight Junctions kann man sich als eine Art Reißverschlussartiges Proteinnetzwerk vorstellen, das sogenannte Verschlusskontakte bildet. Das bedeutet, dass kein Spalt (=kein Interzellularraum) zwischen den benachbarten Zellen übrigbleibt. Die Öffnungen sind also eng, was bereits der Name tight andeutet. Durch diese Abdichtung wird der Transport von Molekülen über die Epithelzellen kontrolliert. Es wird sichergestellt, dass keine Diffusion von gelösten Stoffen zwischen den Epithelzellen möglich ist.

Bei höheren Pflanzen ist hier übrigens analog dazu der Caspary-Streifen zu finden. Dieser sorgt in der Wurzel dafür, dass Schwermetallionen und andere Schadstoffe nicht in die Pflanze gelangen.

Tight Junctions sind aus einem gürtelartigen Netzwerk aus Membranproteinen um die Epithelzellen herum aufgebaut.

Das Netzwerk ergibt sich, indem spezielle Proteine in den Plasmamembranen der jeweiligen Epithelzellen aneinanderbinden. Diese Proteine sind in Strängen angeordnet und bilden eine Aneinanderreihung der einzelnen Kontaktpunkte.

Die wichtigsten daran beteiligten Membranproteine sind Occludin oder Proteine der Gruppe der Claudine. Je nach Zusammensetzung dieser Proteine ergeben sich unterschiedlich dichte Barrieren. Im Darm ist das Netzwerk beispielsweise loser aufgebaut, während es an der Blut-Hirn-Schranke einen dichteren Aufbau aufweist.

Tight Junctions haben neben der mechanischen Funktion, die darin besteht für Stabilität im Epithelverband zu sorgen, noch zwei weitere Funktionen: die Barrierefunktion und die Zaunfunktion.

Durch das komplexe Proteinnetzwerk der Tight Junctions entsteht eine Diffusionsbarriere, die eine Bewegung von gelösten Stoffen durch den Raum zwischen den Epithelzellen verhindert. Dadurch kann der Stoffaustausch auf zellulärer Ebene zwischen Körper und Umgebung kontrolliert werden. So muss beispielsweise jeder Stoff, bevor er vom Darm in den Körper gelangt, die Epithelzellen des Darms durchqueren. Auch in der Blasenwand sorgen Tight Junctions dafür, dass Urin nicht in die Bauchhöhle gelangt.

Außerdem besitzen Tight Junctions eine sogenannte Zaunfunktion. Darunter versteht man, dass die freie Bewegung der Membrankomponenten selbst verhindert wird. Das bedeutet, dass eine seitliche (=laterale) Diffusion der Membranproteine und Phospholipide gemäß dem Flüssig-Mosaik-Modell nicht möglich ist.

Dadurch wird gewährleistet, dass verschiedene funktionelle Bereiche mit unterschiedlicher Membranzusammensetzung in den Epithelzellen bestehen können. Die apikale Region (= Scheitelregion) unterscheidet sich beispielsweise von der basalen Seite (= Unterseite). Dies bezeichnet man als Zellpolarität.

Mikrovilli und Mikrofilamente

Zelloberflächen tierischer Zellen können auch unterschiedliche Ausbuchtungen enthalten. In den Tierzellen sind die wahrscheinlich wichtigsten Membranausbuchtungen die Mikrovilli. An ihrer Ausbildung sind vor allem Aktinfilamente beteiligt.

Die sogenannten Mikrofilamente oder auch Aktinfilamente sind Bestandteile des Cytoskeletts (Abb. 156). Sie sind mit einer Größe von nur etwa 6 nm die kleinsten der drei Filamente.

Mikrofilamente (Aktin)

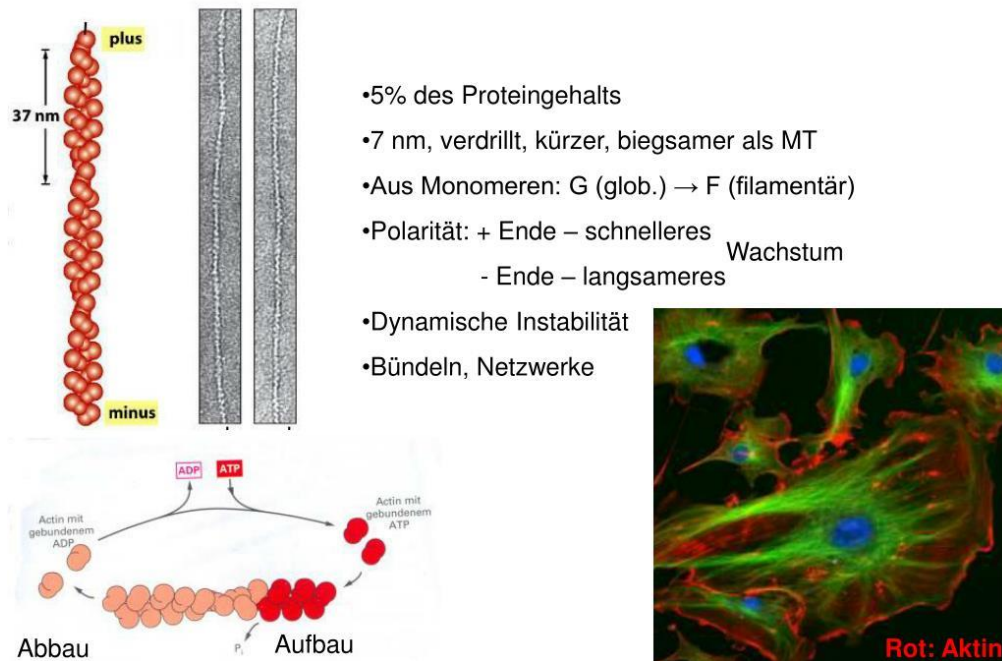


Abbildung 156 Mikrofilamente

Wie der Name Aktinfilament schon sagt, besteht es zum Großteil aus dem Strukturprotein namens Aktin, die die Monomere der Mikrofilamente bilden. Aktin ist eines der häufigsten Proteine in Eukaryoten. Sie entstehen durch Polymerisation aus einzelnen Molekülen des Strukturproteins Aktin.

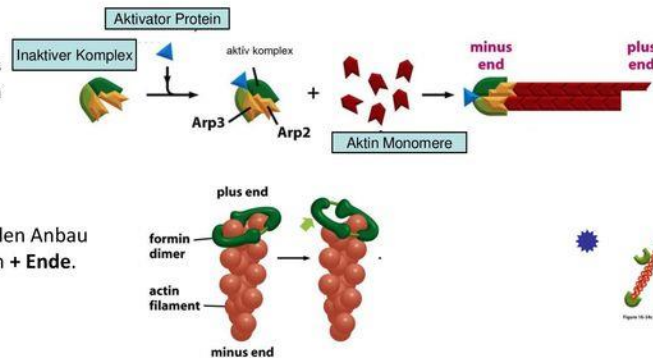
Dabei bindet ein globuläres Aktin an ein ATP. Dieses Monomer (ATP-Aktin) kann sich nun mit weiteren Aktinmolekülen verbinden – polymerisieren, wobei ATP-Aktin unter Abspaltung (Hydrolyse) eines anorganischen Phosphatrestes zu ADP-Aktin wird. Die entstehende Kette von Aktinmonomeren bildet so die filamentöse Form der Aktinfilamente, auch F-Aktin genannt. Aktinfilamente sind aus zwei helikal miteinander verwundenen Aktineinzelfäden aufgebaut. Aktinfilamente weisen eine Polarität auf und haben ein schnell polymerisierendes sogenanntes (+)-Ende und ein langsam polymerisierendes (-)-Ende. ATP-Aktin bindet bevorzugt am (+)-Ende und das Filament wächst an diesem Ende. Am (-)-Ende läuft die Hydrolyse von ATP zu ADP schneller als die Anlagerung eines neuen ATP-Aktins ab, sodass ADP-Aktin dissoziiert und das Filament von dieser Seite verkürzt wird. Zahlreiche Begleitproteine steuern die Polymerisations- und Abbauvorgänge. Im Muskel werden die Filamente beispielsweise durch das Tropomyosin stabilisiert, das sich auf ganzer Länge an ein Filament anlegt. In Zellen außerhalb des Herzens und der Skelettmuskulatur wird Caldesmon gebildet.

Eine große Gruppe von Begleitproteinen, die auch als Aktin-bindende Proteine (ABP) bezeichnet werden, vernetzt Aktinfilamente untereinander und mit anderen Proteinen. Fimbrin, Villin, Filamin und Espin bilden Querverbindungen und so mechanisch steife Bündel (Abb. 157).

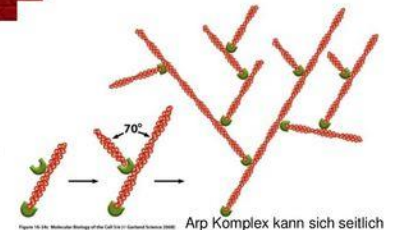
Aktin-assoziierte Proteine

Arp komplex:

Nukleationsstelle für das Auswachsen eines neuen Mikrofilamentes. Bleibt am **– Ende** des Filamentes.



Formin: Hilfsprotein für den Anbau neuer Aktin-Moleküle am **+ Ende**.



Arp Komplex kann sich seitlich an Mikrofilamente binden und dadurch neue Mikrofilamente bilden („Verzweigung“ der Mikrofilamente).

Thymosin: bindet an Aktinmonomere und hemmt die Polymerisierung.

Gelsolin: zerstückelt die Mikrofilamente.

Tropomyosin: stabilisiert die Mikrofilamente durch Bindung seitlich an die Mikrofilamente

Quervernetzende Proteine: quervernetzt die parallelverlaufenden Mikrofilamente in einem Bündel: **α -aktinin, fimbrin**

Filamin: bindet sich an überkreuzende Mikrofilamente in einem Geflecht von Mikrofilamenten (lockere Bindung).

capZ: bindet an das **+** Ende der Mikrofilamente, hemmt dadurch den Auf- und Abbau des Mikrofilamentes, stabilisiert.

ERM Proteine: binden die Mikrofilamente an die **Zellmembran**.

Spectrin: Tetramere aus fibrösen Proteinen, bilden 2D Netze mit Mikrofilamenten (**Membranskelett**).

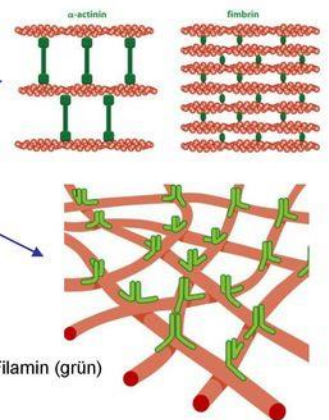


Abbildung 157 Aktin-Assoziierte Proteine.

Aktin und Aktin-bindende Proteine sind an der Ausbildung der Mikrovilli (Abb. 158) beteiligt. Die Mikrovilli kommen vor allem in tierischen Epithelzellen vor.

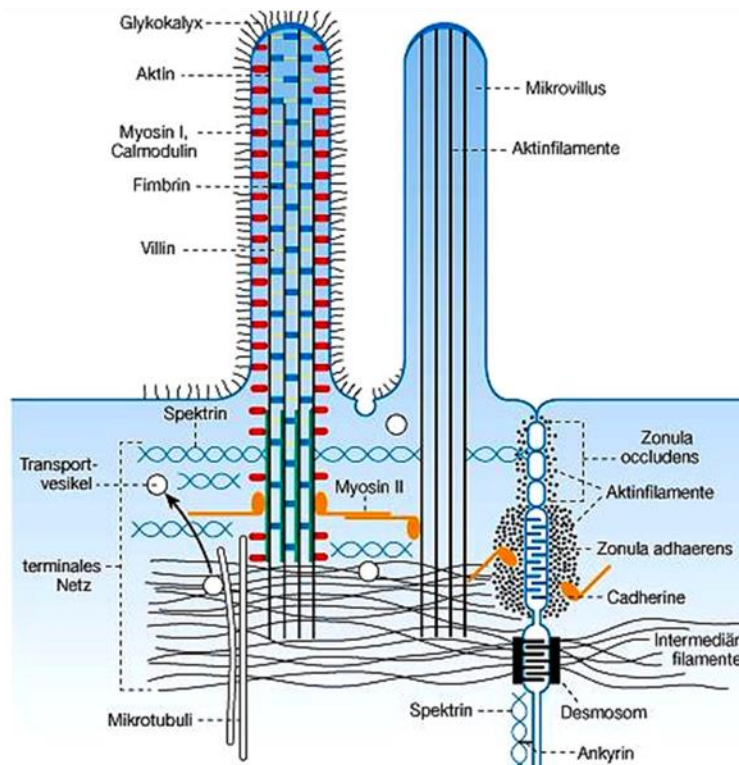


Abbildung 158 Mikrovilli

Mikrovilli sind fadenförmige, leicht bewegliche Ausstülpungen der Zellmembran. Diese sind aus mehreren Bündeln von Aktinfilamenten aufgebaut. Ein Mikrovillus hat in etwa einen Durchmesser von 50 – 100 nm und ist 1 – 2 μm lang. Er besteht aus mehreren Bündeln von Aktinfilamenten. Die Bündel werden über die Proteine Fascin, Fimbrin und Villin in ihrer Form gehalten und stabilisiert.

Ihre Hauptfunktion besteht in der Oberflächenvergrößerung von Zellen. Durch diese größere Oberfläche können die Zellen Stoffe besser aufnehmen und abgeben. Diese ist vor allem bei den Epithelzellen wichtig, die körpereigene oder körperfremde Stoffe aufnehmen (= resorbierende Epithelien). Diese Zellen findet man vor allem in der Darmwand von Dünndarm und Dickdarm oder in den Nieren. Mit einer größeren Oberfläche können diese Stoffe schneller aufgenommen und abgegeben werden. Ohne die Mikrovilli würde der Dünndarm nur eine Oberfläche von etwa 4m^2 besitzen. Mit ihnen kann sich die Fläche des Dünndarms auf bis zu 200m^2 erstrecken.

Die Mikrovilli liegen nur auf der äußeren, der sogenannten apikalen Seite der Epithelzellen. Die innere, basale Seite enthält keine Ausstülpungen.

Um die Mikrovilli zusätzlich zu stabilisieren, ragen die Aktinfilamente mit dem Minus-Ende weiter in die Zelle hinein. Dort bilden sie das sogenannte Terminalgeflecht, das vor allem durch die Strukturproteine Tropomyosin und Myosin zusammengehalten wird. Das Strukturprotein Spektrin verankert die Mikrovilli weiterhin noch am Cytoskelett.

In ihrer Gesamtheit bilden die Mikrovilli den sogenannten Bürstensaum. Dieser Bürstensaum ist von einer Schicht aus Glykoproteinen und Glykolipiden namens Glykocalyx überzogen, die die Zellmembran stabilisiert.

Mikrofilamente und Muskelgewebe

Eine weitere Aufgabe der Aktinfilamente findet man in Muskeln wieder (Abb. 159). Innerhalb der Muskelzellen liegen besonders viele Aktinfilamente. Diese tragen dazu bei, dass die Muskeln anspannen (Kontraktion) oder entspannen (Relaxation) können.

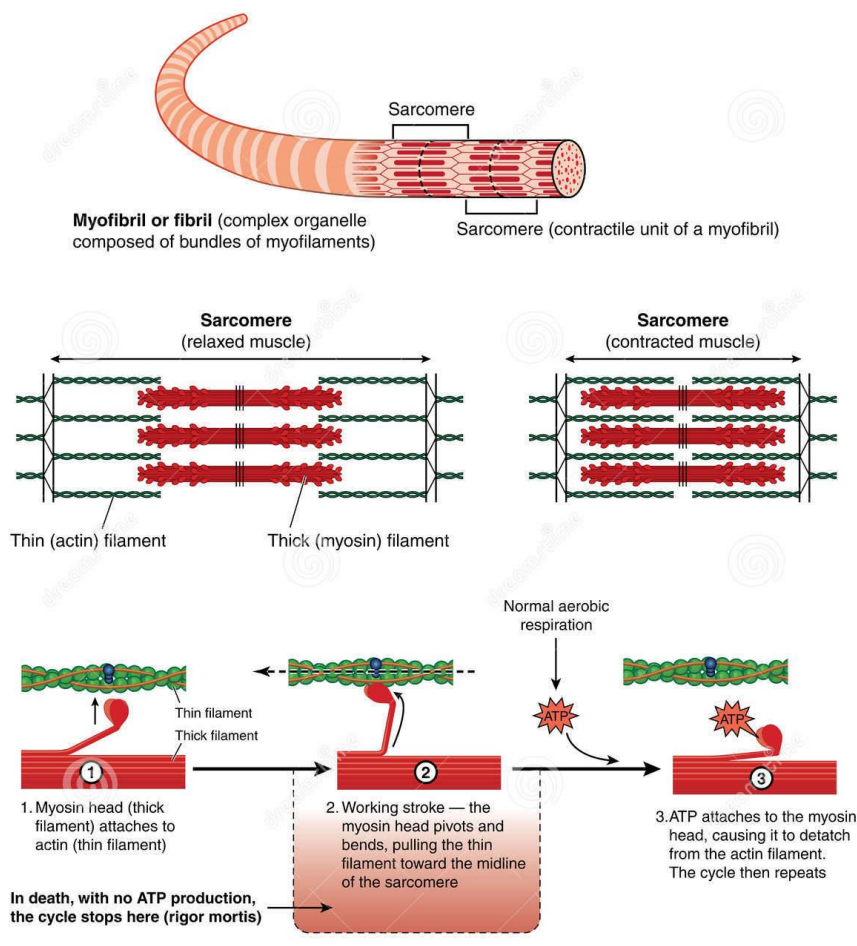


Abbildung 159 Muskelgewebe

Die Muskelkontraktion ist ein biologischer Prozess, bei dem mechanische Kräfte im Muskelgewebe erzeugt werden. Im Falle der Skelettmuskeln werden diese Kräfte durch Sehnen auf die Knochen übertragen.

Die Kräfte entstehen durch Umwandlung von chemischer in mechanische Energie mittels des Aktin-Myosin-Komplexes in den einzelnen Muskelzellen, der wiederum seine chemische Energie aus der Hydrolyse von ATP bezieht.

Diese Bewegung wird durch Änderungen der chemischen Konfiguration und damit der Form der Myosin-Moleküle ermöglicht: Das Myosin besitzt kleine Fortsätze („Köpfe“), die ihren Winkel zum Rest des Moleküls („Schaft“) verändern können. Die Köpfe können wiederum an die Aktin-Filamente binden und diese in sogenannten „Ruderbewegungen“ verschieben. Im Ruhezustand (entspannter Muskel) ist das

Aktinfilament mit so genannten Tropomyosinfäden umschlungen, die die Bindungsstellen der Myosinköpfchen an dem Aktinfilament bedecken. An das Myosin ist ATP gebunden, das Köpfchen befindet sich in einem 90-Grad-Winkel zum Schaft des Moleküls.

Ein Nervenimpuls bewirkt die Ausschüttung von Calcium (Ca^{2+}). Das hat zwei Folgen: Zum einen aktiviert Ca^{2+} die Abspaltung von ATP in ADP + P am Myosinköpfchen, Als Cofaktor wird zusätzlich Mg^{2+} gebraucht. Das Calcium bindet zum anderen an Troponin, das an den Tropomyosinfäden angelagert ist, und verändert dabei deren Konfiguration so, dass die Bindungsstellen freigegeben werden und das Myosin an das Aktin binden kann.

Sobald das Myosin an das Aktin gebunden hat, wird das immer noch am Myosinköpfchen anliegende P_i und kurz danach auch das ADP freigesetzt. Dadurch wird die Verspannung des Myosins in mechanische Energie umgesetzt: Die Myosinköpfchen kippen in einem 45 Grad-Winkel zum Myosinfilament und ziehen dabei die Aktinfilamente von rechts und links zur Sarkomermite.

Der Zyklus wird dadurch abgeschlossen, dass sich neues ATP an das Myosin anlagert. Dadurch löst sich das Myosinköpfchen vom Aktinfilament und die beiden Proteine befinden sich wieder im Ausgangszustand. Steht ATP nicht mehr zur Verfügung, können sich die Moleküle nicht mehr voneinander lösen und es kommt zur Totenstarre.

Pflanzengewebe

Pflanzliche Zellen verfügen über eine Zellwand, die aus Cellulose besteht. Einige Algen und viele Pilze haben anstelle von Cellulose Chitin in ihre Zellwand eingebaut.

Pflanzliche Gewebe lassen sich in zwei Gewebetypen unterscheiden: Bildungsgewebe (Abb. 160) und Dauergewebe (Abb. 161). Bildungsgewebe besteht aus teilungsfähigen Embryonalzellen (dem Meristem).

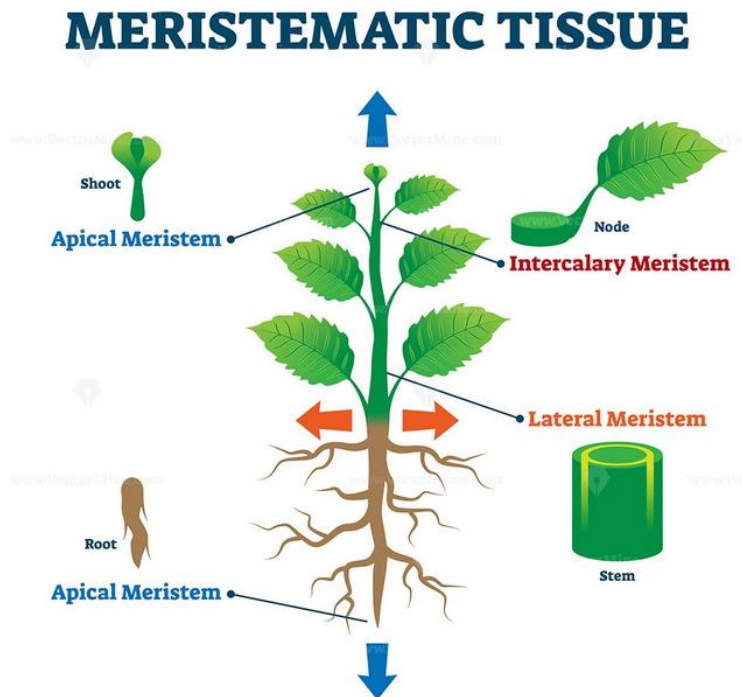
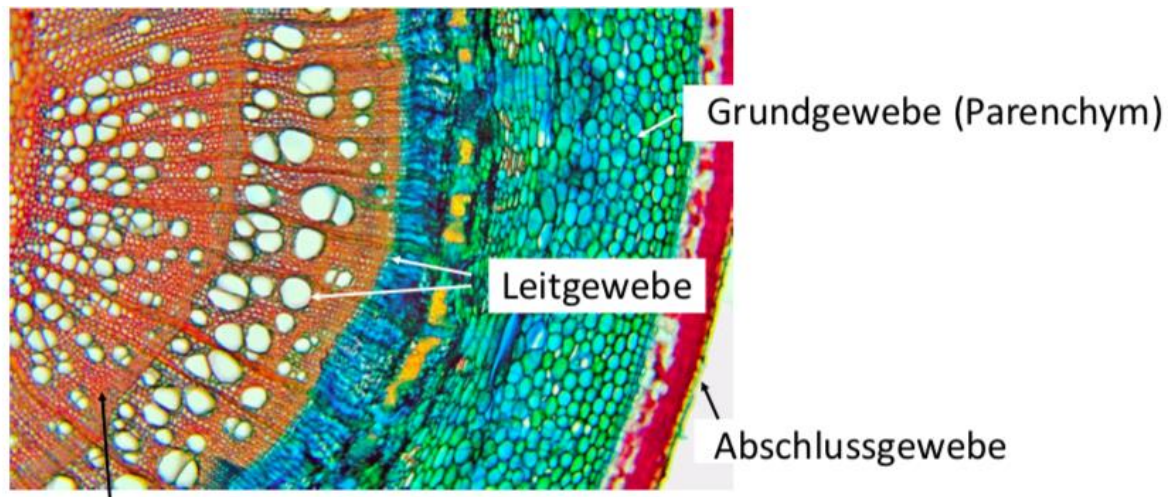


Abbildung 160 Meristem-Gewebe

Die Zellen der Meristeme besitzen in der Regel dünne Zellwände mit wenig Cellulose und sind theoretisch unbegrenzt teilungsfähig. So gibt es Primärmeristeme wie das Protoderm, welches die Epidermiszellen bildet, während z. B. das Prokambium die Leitbündel bildet und das Grundmeristem das Markgewebe. Primärmeristeme bilden sich aus Apikalmeristemen und sind für das Primärwachstum der Pflanze zuständig. Zum Sekundärmeristem gehört das Kambium, eine Gewebeschicht, die für das Dickenwachstum der Pflanzen zuständig ist. Bei Bäumen ist es die hohlzylinderförmige Wachstumsschicht zwischen der Splintholzzone und der Rinde, die auch Kambiumring genannt wird.

Wenn die Zellen nicht mehr teilungsfähig sind, handelt es sich um Dauergewebe (Abb. 161).



Festigungsgewebe

Abbildung 161 Dauergewebe

Bei Pflanzen sind folgende Dauergewebe bekannt.

1. a) Grundgewebe: hierzu zählt das Parenchym. Parenchymzellen sind dünnwandige Zellen des Grundgewebes, die den Großteil von nichtholzartigen (krautigen) Pflanzenstrukturen ausmachen und beispielsweise der Speicherung von Nährstoffen dienen.
2. b) Festigungsgewebe: hier wird in Kollenchym und Sklerenchym unterschieden. Das Kollenchym besteht aus lebenden Zellen, mit nicht verholzten, dehnungsfähigen Zellwänden, das Sklerenchym aus toten Zellen mit verholzten Zellwänden.
3. c) Das Abschlussgewebe welches in Epidermis und Periderm unterteilt wird. Die Epidermis ist das primäre Abschlussgewebe von Sprossachse und Blättern bei höheren Pflanzen. Das Periderm ist das sekundäre Abschlussgewebe der Sprossachse und das tertiäre Abschlussgewebe der Wurzel.
4. d) Das Leitgewebe, welches in Xylem und Phloem geteilt wird. Das Xylem dient dem Wassertransport, das Phloem von Stoffwechselprodukten wie dem Zucker.

Diese vier Gewebetypen bilden bei höher entwickelten Pflanzen (also den sog. Kormophyten wozu die Farngewächse und Samenpflanzen gehören, nicht aber die Moose und Grünalgen) die typischen Organsysteme: Wurzel, Sprossachse und Blätter.

Plasmodesmen

Die Plasmodesmen sind kleine Plasmastränge in Pflanzenzellen (Abb. 162). Sie führen durch Aussparungen (= Tüpfel) in der Zellwand zu einer Nachbarzelle und schaffen so eine Verbindung, bei dem Stoffe transportiert werden können. Ein Plasmodesmos ist ein dünner Strang aus Plasma mit einem Durchmesser von etwa 30-50 nm.

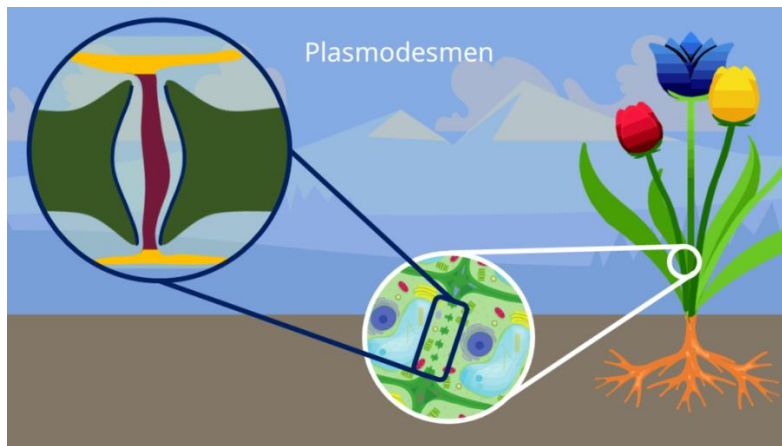


Abbildung 162 Plasmodesmen

Innerhalb der Pflanzenzelle können sich auch mehrere Plasmodesmen befinden.

Da sie während der Zellteilung bestehen bleiben, sind sie in der Lage, die neu gebildeten Tochterzellen miteinander zu verbinden.

Abgesehen von den sogenannten Schließzellen enthalten wahrscheinlich alle Pflanzenzellen Plasmodesmen.

Schließzellen sind bohnenförmige Zellen mit Spaltöffnungen namens Stomata. Ihre Aufgabe ist insbesondere der Gasaustausch mit der umliegenden Luft durch Öffnen und Schließen der Stomata.

In seinem Inneren ist ein Plasmodesmos von einem Strang namens Desmotubulus durchzogen.

Seine Funktionen sind die Stabilisierung des Plasmodesmenkanal und die Verbindung zum Endoplasmatischen Retikulum und den Zellmembranen anderer Zellen. Insbesondere er ist in der Lage, Lipide, Proteine, RNA und weitere Moleküle zur anderen Zelle zu transportieren.

Es gibt drei Möglichkeiten, wie der Desmotubulus die Stoffe transportieren kann:

Zum einen kann er sie entlang seiner Membran bewegen. Zum anderen kann er den Stofftransport durch den zytoplasmatischen Bereich der Plasmodesmen steuern. Eine dritte Möglichkeit besteht darin, die Stoffe durch seinen Innenraum (= Lumen) zu transportieren.

Vor allem bei Letzterem ist er in der Lage, die Größe der weitergeleiteten Stoffe durch seinen Durchmesser zu begrenzen.

Für den Stofftransport innerhalb der Pflanzenzelle existieren zwei unterschiedliche Wege namens Apoplast und Symplast (Abb. 163). Diese beiden Transportwege werden durch die Plasmodesmen jeweils voneinander abgetrennt.

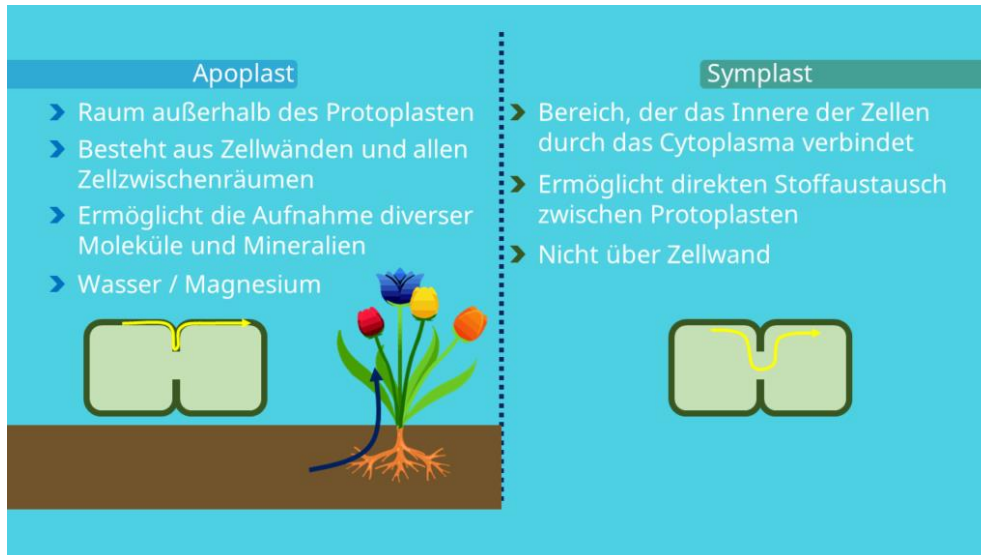


Abbildung 163 Apoplast und Symplast

Der Apoplast bezeichnet den Raum außerhalb des Protoplasten. Dieser Raum besteht aus den Zellwänden und allen Zellzwischenräumen.

Über den Apoplasten ist die Zelle in der Lage, verschiedenste Moleküle oder Mineralien aufzunehmen. Das können zum Beispiel Wasser oder Magnesium sein. Beides erhält die Pflanze vom Boden. Diese Stoffe werden dann über die Zellwände verbreitet.

Der zweite Transportweg geht über den sogenannten Symplasten. Er bezeichnet den Bereich, der das Innere der Zellen durch das Cytoplasma miteinander verbindet.

Über den Symplast können Moleküle und Mineralien direkt zwischen den jeweiligen Protoplasten weitergegeben werden. Diese Weitergabe erfolgt hier nicht über die Zellwand, sondern direkt über die Plasmodesmen.

Kapitel 17: Zelltod und Krebs

Wir wissen wie Zellen sich vermehren. Sie können aber auch sterben.

Zellen eines vielzelligen Lebewesens können auf zwei Weisen absterben. Die erste Art von Zelltod ist die Nekrose (unkontrollierter Zelltod); sie tritt auf, wenn Zellen oder Gewebe durch Gewalteinwirkung, Gift, Sauerstoffmangel oder Nährstoffmangel geschädigt werden. Nekrotische Zellen schwellen häufig an und platzen, sodass ihr Inhalt in den Extrazellularraum freigesetzt wird und eine Entzündung hervorrufen kann.

Die Apoptose (programmierter Zelltod) ist eine genetisch vorgegebene Abfolge von Ereignissen, die im Verlauf normaler Entwicklungsprozesse, aber auch in erwachsenen (adulten) Geweben stattfinden.

Durch den programmierten Zelltod werden Zellen beseitigt, die für den Organismus nicht vorteilhaft sind, weil sie entweder nicht mehr gebraucht werden oder weil Zellen, die zu lange Leben anfällig für Krebs sind. Der Ablauf der Apoptose ist bei vielen Lebewesen gleich. Die Zelle löst sich von ihren Nachbarn ab und ihr Chromatin wird von Enzymen abgebaut, die die DNA in Fragmente von etwa 180 Nucleotiden zerschneiden. Die Zelle bildet membranumhüllte Ausstülpungen („Blasen“) und löst sich in Zellfragmente auf. Eine Reihe verschiedener externer oder interner Signale kann den programmierten Zelltod auslösen (Abb. 164). Diese Signale sind beispielsweise Hormone, Wachstumsfaktoren, die Infektion mit einem Virus, bestimmte Toxine oder eine erhebliche Schädigung der DNA. Diese Signale aktivieren spezifische Rezeptoren, die wiederum Signalwege in Gang setzen, die zur Apoptose führen. Einige apoptotische Signalwege haben die Mitochondrien zum Ziel, etwa indem sie die Durchlässigkeit der mitochondrialen Membranen erhöhen. Die Zelle stirbt schnell ab, wenn die Mitochondrien die zelluläre Atmung nicht mehr ausführen können. Im Verlauf der Apoptose wird eine bestimmte Gruppe von Enzymen aktiviert, die man als Caspasen bezeichnet. Dabei handelt es sich um Proteasen, die in einer Kaskade von Ereignissen Zielmoleküle hydrolysieren. Die Zelle stirbt, wenn die Caspasen die Proteine der Kernhülle, des Cytoskeletts und der Plasmamembran hydrolysieren.

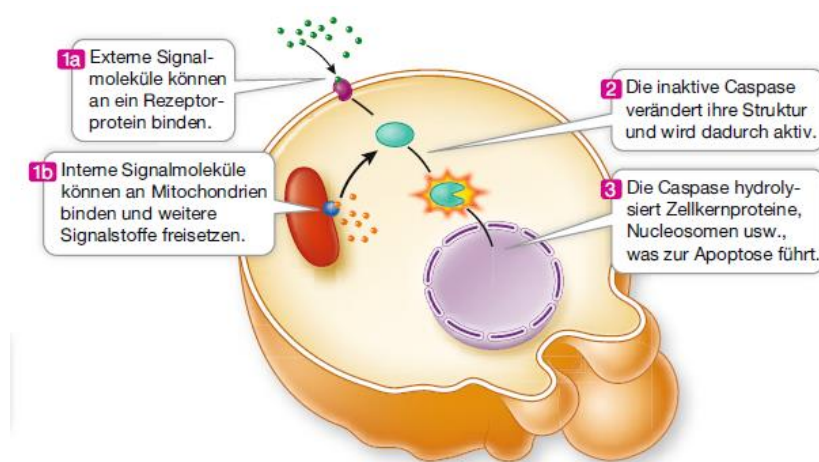


Abbildung 164 Apoptose

Krebs

Durch die Mitose werden einem Lebewesen Zellen hinzugefügt, durch die Apoptose werden sie beseitigt. Unter normalen Bedingungen befinden sich diese Vorgänge im Gleichgewicht, sodass die Situation für den Organismus insgesamt von Vorteil ist. In anderen Situationen führt sie zu einem unkontrollierten Zellwachstum, geläufig auch als Krebs bezeichnet. Eine Krankheit, die den meisten ernste Sorgen bereitet und diese sind berechtigt. Einer von drei Bewohnern der USA bekommt während seines Lebens irgendeine Form von Krebs und zurzeit stirbt daran jeder Vierte. Mit 1,5 Mio. Neuerkrankungen und 0,5Mio. Todesfällen pro Jahr liegt Krebs in den USA als Todesursache auf Platz zwei hinter den Herz-Kreislauf-Erkrankungen. In Deutschland sind es rund 340.000 Neuerkrankungen und 210.000 Krebstote pro Jahr.

Krebszellen (böartige Tumorzellen) unterscheiden sich von den normalen Zellen, aus denen sie hervorgehen, durch zwei Eigenschaften: Krebszellen verlieren die Kontrolle über die Zellteilung. Krebszellen können in andere Körperregionen wandern.

Die meisten Zellen im Körper teilen sich nur, wenn sie extrazelluläre Signale erhalten, beispielsweise Wachstumsfaktoren.

Krebszellen reagieren nicht auf diese Kontrollen und teilen sich stattdessen mehr oder weniger kontinuierlich. Schließlich bilden sie Tumoren (große Massen von Zellen). Tumoren können gutartig oder böartig sein. Gutartige Tumoren ähneln dem Gewebe, aus dem sie stammen. Sie wachsen langsam und bleiben auf die Region beschränkt, in der sie sich entwickeln. Gutartige Tumoren sind kein Krebs, aber sie müssen entfernt werden, wenn sie so auf ein Organ übergreifen, dass sie seine Funktion beeinträchtigen.

Bösartige Tumoren sehen überhaupt nicht aus wie das ursprüngliche Gewebe. Bösartige Tumorzellen haben häufig irreguläre Strukturen, etwa Zellkerne von unterschiedlicher Größe und Form. Die zweite und besonders bedrohliche Eigenschaft von Krebszellen ist ihre Fähigkeit, in das umgebende Gewebe einzudringen und sich in andere Körperregionen auszubreiten, indem sie sich durch das Blut oder die Lymphe transportieren lassen. Haben sich böartige Tumorzellen in irgendeiner weiter entfernten Körperregion angesiedelt, teilen sie sich und bilden so an dieser Stelle einen neuen Tumor, als Tochtergeschwulst oder Metastase bezeichnet. Diese Ausbreitung führt an verschiedenen Stellen des Körpers zu Organversagen und erschwert die Krebsbekämpfung in besonderer Weise.

Wir haben im Kapitel zur Mitose und Zellzyklus kennengelernt, dass bestimmte Signale das Wachstum und die Teilung von Zellen hemmen oder fördern. Wachstumshormone fördern die Zellteilung, andere Proteine wie das Retinoblastomprotein (RB) hemmen den Zellzyklus. Krebszellen reagieren auf diese Signale nicht. Hier sind zwei entsprechend wichtige Gengruppen wichtig: Onkogene und Tumorsuppressorgene (Abb. 165).

Onkogenproteine werden von Onkogenen exprimiert und sind positive Regulatoren von Krebszellen. Sie leiten sich von normalen positiven Regulatoren ab und sind entweder mutiert, sodass sie eine übermäßige Aktivität zeigen, oder sie kommen im Überschuss vor. So können sie Krebszellen dazu stimulieren, sich häufiger zu teilen.

Die Zielstrukturen von Onkogenproteinen können Wachstumsfaktoren, deren Rezeptoren oder andere Komponenten eines Signaltransduktionswegs sein, der die Zellteilung stimuliert.

Ein Beispiel für ein Onkogenprotein ist ein Wachstumsfaktorrezeptor bei Brustkrebszellen. Normale Brustzellen enthalten relativ geringe Mengen des Wachstumsfaktorrezeptors HER2. Deshalb werden Brustepithelzellen normalerweise bei Vorhandensein dieses Wachstumsfaktors auch nicht dazu stimuliert, sich zu vermehren. Bei etwa 25% aller Fälle von Brustkrebs resultiert eine Mutation der DNA in einer erhöhten Produktion des HER2-Rezeptors. Das führt zu einer positiven Stimulation des Zellzyklus und zu einer schnellen Proliferation der Zellen mit der mutierten DNA.

Tumorsuppressorproteine werden von Tumorsuppressorgenen exprimiert und kommen sowohl in Krebszellen als auch in normalen Zellen vor. In normalen Zellen wirken sie als negative Regulatoren, in Krebszellen sind sie inaktiv.

Ein Beispiel ist das bereits erwähnte Retinoblastomprotein (RB), das in der G1-Phase als Restriktionspunkt wirkt. In einigen Krebszellen ist RB inaktiv, sodass der Zellzyklus ungebremst voranschreitet. Einige Virusproteine können Tumorsuppressorproteine inaktivieren. So infiziert beispielsweise das menschliche Papillomvirus Epithelzellen im Gebärmutterhals. Ein Genprodukt dieser Viren ist das Protein E7. Es bindet an das RB-Protein und verhindert dadurch, dass es den Zellzyklus hemmt. Ein weiterer wichtiger Tumorsuppressor ist p53, ein Transkriptionsfaktor, der bei Kontrollpunktsignalwegen des Zellzyklus und der Apoptose eine Rolle spielt. Die Bedeutung von p53 als Tumorsuppressor lässt sich dadurch veranschaulichen, dass bei über 50% aller Tumoren des Menschen das p53-codierende Gen mutiert ist.

Da es Onkogene und Tumorsuppressorgene gibt, ist mehr als ein mutiertes Protein erforderlich, damit der Zellzyklus von Krebszellen voranschreiten kann. An einem einzigen Tumor können mehrere mutierte Onkogene und Tumorsuppressorgene beteiligt sein. Bei der Maus gibt es beispielsweise zwei wichtige Onkogene: Myc, dessen Expression den Zellzyklus stimuliert und die Apoptose verhindert, und Ras, ein Signalprotein. In Experimenten ließ sich zeigen, dass beide Onkogene exprimiert werden müssen, damit der Zellzyklus in Mauszellen in Gang gesetzt wird, die dadurch zu Tumorzellen werden.

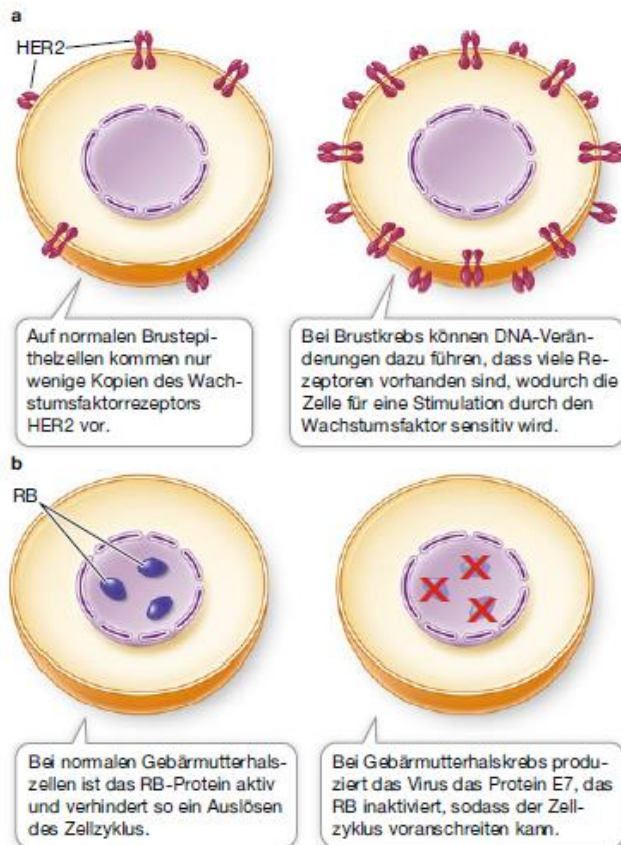


Abbildung 165 Wirkung von Onkogenen (a) und Tumorsuppressorgenen (b).

In normalen, nichtwachsenden Geweben ist die Zellteilungsrate gleich der Apoptoserate, sodass die Zellpopulation insgesamt nicht zunimmt. Bei Krebszellen ist die Regulation des Zellzyklus gestört, sodass sich die Zellteilungsrate erhöht.

Die immer noch erfolgreichste und am häufigsten angewendete Behandlungsmethode gegen Krebs ist ein chirurgischer Eingriff. Das physikalische Entfernen eines Tumors ist zwar effizient, oftmals ist es jedoch für einen Chirurgen schwierig, alle Tumorzellen zu beseitigen. (Ein Tumor von etwa 1 cm Größe enthält bereits 1Mrd. Zellen.) Tumoren sind im Allgemeinen von normalen Geweben umgeben. Darüber hinaus besteht die Möglichkeit, dass sich Zellen vom Tumor abgelöst und in andere Organe ausgebreitet haben. So ist es unwahrscheinlich, dass ein lokal begrenzter chirurgischer Eingriff eine dauerhafte Heilung bewirkt. Deshalb gibt es zusätzliche Verfahren, um Krebs zu heilen, und diese sind vor allem auf den Zellzyklus gerichtet. Das Ziel ist, bei den Tumorzellen die Zellteilungsrate zu erniedrigen oder die Apoptoserate zu erhöhen, sodass ihre Population abnimmt.

Ein Beispiel für einen Wirkstoff gegen Krebs, der sich gegen den Zellzyklus richtet, ist Fluoruracil, das die Synthese von Thymin blockiert, eine der vier Basen in der DNA. Der Wirkstoff Paclitaxel (im Handel z. B. unter dem Namen Taxol) behindert die Funktion der Mikrotubuli in der Mitosespindel.

Beide Wirkstoffe hemmen den Zellzyklus, und die Apoptose bewirkt, dass der Tumor schrumpft. Drastischere Auswirkungen zeigt die Strahlentherapie, bei der hochenergetische Strahlung auf den Tumor gerichtet wird. Dadurch kommt es zu

erheblichen DNA-Schäden, und der Kontrollpunkt für die DNA-Reparatur wird überlastet. Das wiederum löst bei den betroffenen Zellen die Apoptose aus. Ein ähnliches Ziel verfolgt die Chemotherapie, bei der versucht wird, die Krebszellen gezielt zu vergiften.

Ein Hauptproblem bei diesen Behandlungsmethoden besteht darin, dass neben den Tumorzellen auch normale Zellen angegriffen werden. Diese Behandlungsmethoden sind insbesondere für solche Gewebe schädlich, die im Normalzustand eine hohe Zellvermehrungsrate haben, wie das Darmepithel, die Oberhaut (Epidermis) oder das Knochenmark (das laufend Blutzellen produziert).

Das vorherrschende Ziel in der Krebsforschung ist die Entwicklung von Behandlungsmethoden, die spezifisch nur Krebszellen angreifen. Ein vielversprechendes Beispiel ist Herceptin, das gegen den HER2-Wachstumsfaktorrezeptor gerichtet ist, der wiederum an der Oberfläche einiger Typen von Brustkrebszellen auf hohem Niveau exprimiert wird. Herceptin bindet spezifisch an den HER2-Rezeptor, stimuliert ihn jedoch nicht. Das verhindert ein Andocken des natürlichen Wachstumsfaktors, sodass die Zellteilung nicht angeregt wird. Da die Zellteilung blockiert ist und die Apoptoserate gleichbleibt, schrumpft der Tumor.

Einen ganz anderen und in den letzten Jahren zunehmend erfolgversprechenden Behandlungsansatz von Krebs verfolgt die Immuntherapie. Dabei wird versucht, das Immunsystem des Patienten zur Bekämpfung der Krebszellen anzuregen.