

Diese Ausarbeitung begann als Entwurf für

Verstärker ganz allgemein

Der ständig erweitert und verbessert werden sollte. Während mein altes Buch

Das Mensch - Technik - System. Expert-Verlag, Renningen-Malmsheim 1999 + Linde - Verlag

die mögliche Verkopplung beider Gebiete behandelt, wurde hier zunächst versucht, Vergrößerung, Erweiterung und Steigerung der menschlichen Leistungen, vor allem der Sinne und Muskeln herauszuarbeiten. Das trifft besonders die Kräfte, das Sehen und Hören. Dabei sollen möglichst alle denkbaren Varianten von Verstärkung erfasst und verglichen werden. Als eine erste Grundlage dafür gab es bereits die drei Vorträge bzw. Folien:

- 09.04.11 *Verstärker.pdf*. Der erste Versuch das Prinzip Verstärker zu erklären
- 29.04.11 *Wirkungen.pdf*: Wie Wirkungen zustande kommen, Prinzip Verstärkung, Verbessert und erweitert von einer älteren Fassung
- 12.12.11 *infVerstärk.pdf*: Zusammenhang und Erklärung von allgemeiner Verstärkung und Information

Teilweise gehört dazu auch meine **Vorlesung** mit dem Ziel: Alles was sichtbar gemacht werden kann. Die dazugehörige Skripte sind:

23.02.16 *Visualisierung.pdf*. – 02.03.16 *Visualisierung2.pdf* – 05.03.16 *Visualisierung3.pdf*

10.03.16 *Visualisierung4.pdf* – 09.11.16 *Visualisierung5.pdf*

Eventuell wird sollte so schließlich ein entsprechendes Buch entstehen. Vor allem werden dabei ergänzende Texte, Formeln, Bilder, Fakten entstehen und sich so ein neuer Zusammenhang ergeben. Wichtig ist mir die immer die bildliche Darstellung der Zusammenhänge. Ein zusätzliches Teilziel ist die Anpassung an unsere Sinne und Kräfte bzgl. Raum, Realität, usw. *Ursprünglich* wollte ich dabei den Inhalt nach den Mechanismen aller Verstärkungen bearbeiten und ordnen. Das liegt immer noch aktuell weitgehend vor und wird vielleicht auch bis zu Ende geführt, jedoch nicht mehr ganz so konsequent. Doch abschließend ist – wie sich erst jetzt zeigte – eine Ordnung gemäß der Evolution besser geeignet. Das bedingt auch eine Änderung des Titels, in etwa:

Evolution von Verstärkung

durch vergrößern, verstärken, abschwächen, beschleunigen und vermehren von Etwas,

sowie Störungen mindern bis unterdrücken und zwar positiv ↔ negativ

Die Stufen sind etwa, s. mein Buch: Speicher als Grundlage für Alles, Shaker Verlag 2019

- *Atomar, Physikalisch*
- *Chemisch*: Entwicklerflüssigkeit, Katalysator. Kontrastmittel
- *Leben*, Biosynthese, Hormon, Antikörper, Immunsystem
- *Sinne, Muskeln*, Geschmacksverstärker, Gewürze, Parfum, Kosmetik Desodorierung
- *Verhalten* Reflex, Konditionierung, Droge usw.
- *Gedächtnis*: Lernen, Motivierung, Hypnose
- *Sozial*: Bestechung, Politik, Macht, Medienwirkungsforschung
- *Technik*: Kraft, Energie, Rechentechnik, Virtualität.

Noch sinnvoller könnte es sein, auch die Speicherung einzuordnen. In beiden Fällen liegen nämlich fast identischen Etappen vor. Dann entstünde zusammengefasst etwa (wird noch erprobt):

Speichern und Verändern (Verstärken) als universelle Grundlagen von Evolution

Damit wird das hier vorliegende Material vorläufig und könnte eventuell zu einer neuen Qualität führen. Ein möglicher Titel wäre auch:

Kybernetik – Information – Realität

Sie könnte sich zeigen, dass wir Alles über unsere Sinne und Handlungen gemäß Information erfahren, erleben. Dabei erweist sich das Prinzip der Kybernetik als grundlegend.

Auf alle Fälle wird hier in Abständen immer der aktuelle Stand wiedergegeben.

1. Einführung

Hier besteht das Ziel, möglichst alle Verstärkerarten zu erfassen und entsprechend dem Vorwort vorrangig zu erklären und nicht nur zu beschreiben. Daher sei zunächst der Unterschied zwischen beiden deutlich herausgestellt.

Beschreiben	Erklären ≈ Begründen
Erfolgt sprachlich, schriftlich, bildlich, mimisch, gestisch. Dient der Informationsweitergabe. Betrifft Erlebnisse, Tatbestände, Personen, Sachen, Ideen. Erfolgt phänomenologisch, deskriptiv. Üblich sind Schilderung, Erzählung, Mitteilung.	Ableiten aus Gesetzen, Theorien, Hypothesen, Zusammenhängen. Abbilden eines Sachverhaltes in einer Theorie. Deuten und Begründen der Zusammenhänge. Logisches Denken, Verstehen, Lehren. Axiomatiken. Zusammenhang von Ursache → Wirkung
Beschreibungs-Sprachen der Rechentechnik für Programme und der Bedienung von Geräten. Es geschieht dies, danach geschieht usw. Wird vorwiegend auswendig als Wissen gelernt	Mitteilung = (offizielle) Äußerung, z. B. Kriegs-, Liebes-, eidesstattliche oder feierliche Erklärung. Mit hoher Verkürzung sind Naturgesetze und Axiomatiken wichtig. Verlangt Intelligenz mit Nutzung von Kombinationen.

Eine **Erklärung** und **Begründung** zum **Wesen** bzw. **Zusammenhang** der Sachverhalte (Wirkungen) vordringen, Gegensatz zur **Beschreibung**. Das ermöglicht u. a. auch **wissenschaftlich Voraussagen**. Dazu müssen die Eigenschaften und das Verhalten aus **Gesetzen, Theorien oder Hypothesen** der Wissenschaft (Physik, Chemie usw.) abgeleitet werden. Bezüglich „Erklärung“ sind zwei Autoren wichtig:

- IMMANUEL KANT (1724 - 1804): „Ableitung aus einem deutlich angebbaren Prinzip“
- GEORG WILHELM FRIEDRICH HEGEL (1770 - 1831): „Rückführung von Erscheinungen auf vertraute Verstandesbestimmungen“

Achtung! Andere Bedeutung bei: *Es spottet jeder Beschreibung.*

Hier wird die Erklärung der Welt mit der Ständigkeit von Objekten, Gesetzen und Konstanten begonnen. U.a. treten dann als Folge der Gesetze Veränderungen auf, wobei Objekte ihren Ort oder/und ihre Eigenschaften ändern. Das kann unterschiedlich schnell und/oder intensiv erfolgen. Damit in der Evolution neue Qualitäten entstehen sind oft intensive Änderungen notwendig. Daher ist es sinnvoll, die Inhalte der heutigen Begriffe stark, Stärke usw. möglichst übersichtlich zu erfassen:

stark

- als Adjektiv: bzgl. Kraft, Ausprägung und Macht auf, und gilt dann ähnlich wie dick, gehaltvoll und großartig.
- als Adverb steht es für viel, reichlich, heftig und sehr.

Stärke

- bestimmt eine große Kraft oder Macht und ist dann zählbar bzgl. Ausprägung, Umfang, Durchmesser, Dicke, Heftigkeit und Intensität
- Sie betrifft auch Konzentration oder Gehalt an gelöstem Stoff, Grad der Leistungsfähigkeit, Bestand (an Menschen)
- Si ist auch eine positive Eigenschaft, besondere Fähigkeit oder ein Vorzug
- Sie ist ein in den Chloroplasten der Pflanzen gebildetes, quellfähiges Polysaccharid, das u. a. zur Versteifung von Textilien benutzt wird.

stärken

Bedeutet: stark machen, erfrischen, erquicken oder mit Pflanzenstärke steifen

Verstärkung, verstärken:

Primär betreffen sie heute elektrische Schaltungen bzgl. Signale und Störungen. Dabei ist der Zusammenhang von **Ursache** ⇒ **Wirkung** wichtig. Die Änderungen erfolgen **zeitlich** und/oder **räumlich**. Beide Begriffe werden heute vielfältig benutzt, z. B. bei der Erhöhung von **Bildkontrasten** und zwar bei fotografischen Negativen als auch bei Röntgenbildern benutzt. Darüber hinaus gibt es auch Geschmacksverstärker und vieles mehr. Daher ist es vorteilhaft, alles ähnliche Geschehen als Verstärkung (verstärken) zu bezeichnen. Dann ergeben sich die **9 Varianten** der folgenden Tabelle. Sie sind mit alphabetisch geordneten Buchstaben bezeichnet, die sowohl funktionell als auch durch das Verstärkte bestimmt sind. Später werden die Varianten einzeln und bevorzugt nach ihrer Entstehungszeit geordnet ausführlicher behandelt. Verwandt mit den Verstärken sind **Oszillatoren** und **Regler**, die im Abschnitt ##### behandelt werden.

Mögliche Arten zur Verstärkung

Verstärker	Betrifft	Beispiele	Ursache - Wirkung
D-	Digital	Schalter, Relais, Taster, Pulslänge, Digitaltechnik	Ja-Nein, Wahr-falsch
E-	Effekte	Laser, Maser, Laufzeit von Teilchen, Stimulierte Emission, Sekundärelektronen	1) viele Varianten
G-	Gesetze, Formeln	Hebel, Schraube, Getriebe $\text{Energie} = \text{Weg} \cdot \text{Kraft}$, $N = U \cdot I$, $U = I \cdot R$	Erhaltungssätze mittels Produkt (Summe)
I-	Impulse	Ramme, Hammer, Axt $\text{Energie} = \text{Leistung} \cdot \text{Zeit}$	Integrierte Kraftstöße
K-	Kybernetik	Katastrophen, Irregularitäten, Resonanz, Mimose, (Auslösung)	2) Blockierte Energie, Steuerung - Regelung, Genetik: Schlüssel-Schloss
O-	Optisch	Lupe, Hohlspiegel, Mikroskop, Fernrohr Lichtablenkung, Bildverstärker, chemisch, auch bei Röntgen und Schall.	primär 2D-Bilder, auch 3D, sowie Schall und Antennen.
P-	Physiologie, Psychologie	Geschmack, Konditionierung, Medienwirkung	3) Reiz-Verhalten
R-	Rekursion	L-Systeme mit Bilderzeugung z. B. durch eine Funktion	sehr häufige eingeschachtelte Wiederholung
S-	Signale	Röhren, Transistoren	Elektrische Felder

1) Maser, Laser, Mikrowellen-, Laufzeit-Röhren (Klystron, Wanderwellenröhre, Magnetron, Scheiben-Triode, Gunn-Element), Parametrische = Reaktanz-Verstärker, Elektronen-Vervielfacher usw.

2) **Genetik** mit Enzymen als Schlüssel-Schloss-Prinzip: der Prozess läuft nur ab, wenn das **Enzym**, der **Katalysator** usw. vorliegt.

3) **Psychologie**: ein Reiz, der als Konsequenz eines bestimmten Verhaltens wirkt. Z. B. **Konditionierung** positiv durch Belohnung, Motivierung, wie Geld, Lob oder negativ durch unangenehme Reize, wie Schmerz oder Hunger. Sie ist wichtig für **Lernen**, Lerntheorien. Auch **Medienwirkung**: Massenkommunikation kann Verstärkung bestehender Einstellungen bewirken. Ferner **Intelligenz-Verstärker** durch Drogen und Computer.

2. G-Verstärker = Austausch von Parametern

G-Verstärker sind wohl die ältesten Anwendungen von Verstärkern. Sie wurden spätestens im alten Griechenland genutzt. Am meisten bekannt sind hier mehrere Nutzenanwendungen von Archimedes (284 - 212 v.Chr.). 250 v.Chr. findet er die Gesetze des Hebels und der schiefen Ebene; 239 v.Chr. den Flaschenzug und wendet ihn in der Hafenstadt Syrakus (Sizilien) an. Mit ihm konnte König Hieron ein großes Schiff allein ins Wasser gleiten lassen. So überzeugte er ihn von der Bedeutung technischer Wissenschaft. 212 v.Chr. setzte er angeblich mit einem riesigen Spiegel die römische Flotte bei Syrakus in Brand. Von ihm soll auch die Aussage stammen, gebt mir einen festen Punkt in der Welt und ich hebe die Erde aus den Angeln (s. Bild 1). Darüber hinaus sind über ihn zwei Anekdoten bedeutsam:

Der Legende nach besaß König Hieron II von Syrakus einen goldenen Kranz (Krone). Er hatte aber den Verdacht, dass er evtl. mit Silber gefälscht worden sei. Daher übertrug er die für damalige Zeit sehr schwierige Prüfung Archimedes. Dieser fand die Problemlösung, als er zum Bade in die Wanne stieg und dabei das Wasser überlief. Dadurch begriff er, so das spezifische Gewicht und damit der Silberanteil zu bestimmen. Die Krone war wirklich gefälscht!

Im 2. Punischen Krieges belagerten die Römer von Syrakus und rückten schließlich 212 v.Chr. unter Leitung des Konsuls Claudius Marcellus als Sieger ein. Dabei hatte er seinen Soldaten ausdrücklich en Befehl erteilt, Archimedes unbedingt zu verschonen. Sie trafen ihn, als er geometrische Figuren für eine Berechnung in den Sand zeichnete. Er war so tief in sein Problem versunken, dass er einen sich nähernden römischen Soldaten aufforderte: „Störe mir meine Kreise nicht!“. Darauf hin habe der ihn erschlagen.

Eine deutlich andere Meinung vertritt hierzu der Tschechische Schriftsteller Karel Čapek (1890 - 1938) in seiner Kurz-Erzählung „Archimedes' Tod“ von 1938. („Wie in alten Zeiten“. Aufbau-Verlag, Berlin 1977, S. 40): Es war den Römern gut bekannt, dass Archimedes viele wirksame Kriegsmaschinen erfunden hatte, um den Feind abzuwehren, u. a. Steinschleudern, Flaschenzüge und Haken, mit denen römische Schiffe hochgezogen werden konnten, um sie dann zerbersten zu lassen, sowie optische Vorrichtungen, um sie in Flammen zu setzen. Deshalb versuchte der Konsul ihn für Rom gewinnen, um mit ihm Roms Weltmacht weiter zu stärken. Doch hierzu war Archimedes nicht zu überzeugen und daher musste er sterben!



Bild 1. Archimedes hebt die Erde aus ihren Angeln.

Hebel, schiefe Ebene, Keil und Schraube

Diese Methoden beruhen alle auf physikalischen Gesetzen, die mehrere physikalische Größen multiplikativ zur einer anderen Größe verbinden. Ein Beispiel ist die Energie E , die sich z. B. aus Weg x und der Kraft F gemäß $E = x \cdot F$ berechnet. Bei konstant bleibender Energie E gemäß dem Energieerhaltungssatz kann daher durch einen passenden Mechanismus die Kraft mittels eines längeren Weges erhöht werden. Das leistet besonders übersichtlich der Hebel (Bild 2a). Das hat bereits Archimedes mehrfach erfolgreich genutzt, und es führte auch zu seinem rein theoretischen Gedanken des Bildes 1.

Der Hebel in Bild 2a ist links drehbar gelagert. Dann ist der Weg x_1 am rechten Ende l_1/l_2 -mal größer als x_2 im Angriffspunkt der zu nutzenden Kraft F_2 . Die dafür aufzuwendende Kraft F_1 am Hebelende ist entsprechend l_2/l_1 mal kleiner.

Neben den einseitigen Hebel (a) existieren noch mehrere Abwandlungen, z. B. der zweiseitige Hebel (b), der auch im Bild 1 vorliegt. Außerdem muss der Hebel nicht gerade sein, sondern kann auch geknickt gestaltet werden.

Ähnlich wirkt sich auch eine **schiefe Ebene** aus (b). Die Fallgeschwindigkeit der Erde G wird durch sie in die Normalkraft zur geneigten Schiefe $F_N = G \cdot \sin(\alpha)$ und in die Kraft zur Beschleunigung der Kugel $F_H = G \cdot \cos(\alpha)$ gewandelt. So verlangsamte 1859 Gallilei für seine Experimente zur Fallgeschwindigkeit das Fallen und erhielt besser messbare Zeiten.

Auch die Wirkung eines **Keiles** (d) kann über Hebel beschrieben werden. Entsprechend dem Keilwinkel 2α verstärkt sich die senkrechte Druckkraft F_1 auf die seitlich wirkende Kraft auf $F_2 = F_1 / (2 \sin(\alpha))$ und treibt so den Block nach beiden Seiten auseinander. Alle scharfen Messer – und ganz besonders die Skalpelle der Mediziner – trennen auf die Weise.

Schließlich sei noch die **Schraube** angeführt. Mit einer Drehung von 180° wird sie um die Ganghöhe h vorangetrieben. An ihrem Rand wird dabei ein Weg von $b = 2\pi r$ benötigt. Eine zusätzliche Verstärkung um R/r erfolgt durch den Außenrand der Schraube.

Ergänzend sei noch erwähnt, dass auch die **reziproke „Übersetzung“** möglich ist. Dann wird zwar weniger Kraft, dafür aber ein längerer Weg verfügbar. Darauf wird noch bei den Getrieben genauer eingegangen.

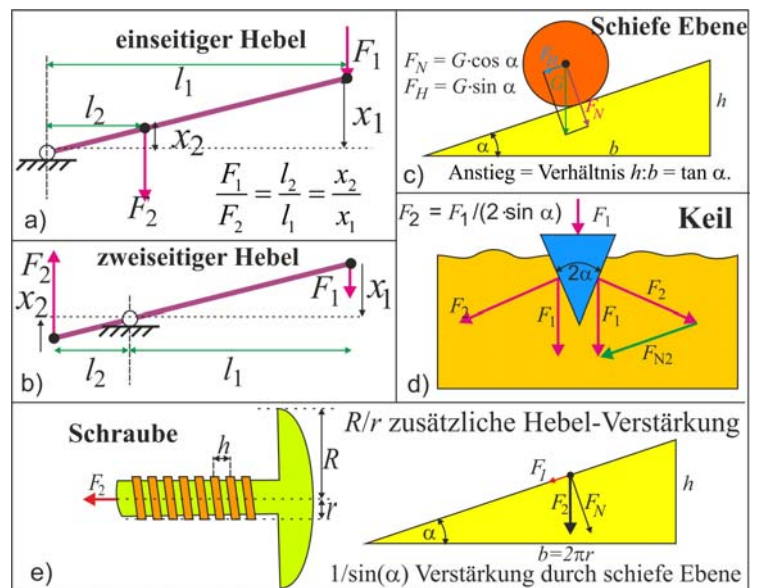


Bild 2. Prinzip von Hebel, schiefe Ebene, Keil und Schraube.

Rollen und Hydraulik

Die oben benutzten Produktbildungen $E = x \cdot F$ sind nur einige Beispiele für die Verstärkung von Kraft. Sie lassen sich leicht auf mehrere Rollen- anordnungen und die Hydraulik übertragen. **Bild 3a** zeigt ein Riemengetriebe, das etwa genauso wie der einseitige Hebel von Bild 2a wirkt. Die Rollendurchmesser r_1 und r_2 entsprechen dabei im Wesentlichen den Teillängen l_1 und l_2 . Die feste Rolle (b) entspricht dem zweiseitigen Hebel (Bild 2b) mit gleichlangen Teilabschnitten. Es erfolgt daher keine Verstärkung, sondern nur eine Richtungsumkehr für die Kraft. Ist die Rolle dagegen beweglich wie in (c), so wird nur noch die halbe Kraft zum Heben benötigt. Besonders leistungsfähig ist der Flaschenzug mit mehreren festen und beweglichen Rollen. Bei n (ganzzahlig) Rollen ist dann neben der Richtungsumkehr für die Kraft F_1 nur $1/n$ das zu hebende Gewicht F_2 erforderlich.

Eine deutlich andere Methode ermöglicht die Hydraulik nach Bild 3e. Sie setzt jedoch inkompressible Flüssigkeit voraus. Dadurch ist der Druck in beiden Zylinder immer gleich groß. Infolge der unterschiedlichen Durchmesser D_1 und D_2 verhalten sich aber die Drucke auf die Kolben wir deren Oberflächen und es gilt $F_2 = F_1 \cdot (D_2/D_1)^2$. Für große Kolbenhübe bei D_2 müsste der Kolben bei D_1 aber sehr lang sein. Das lässt sich aber weitgehend durch Pumpen mit den beiden Ventilen weitgehend vermeiden. So werden heute z. B. riesige Baukräne in der Höhe verschoben.

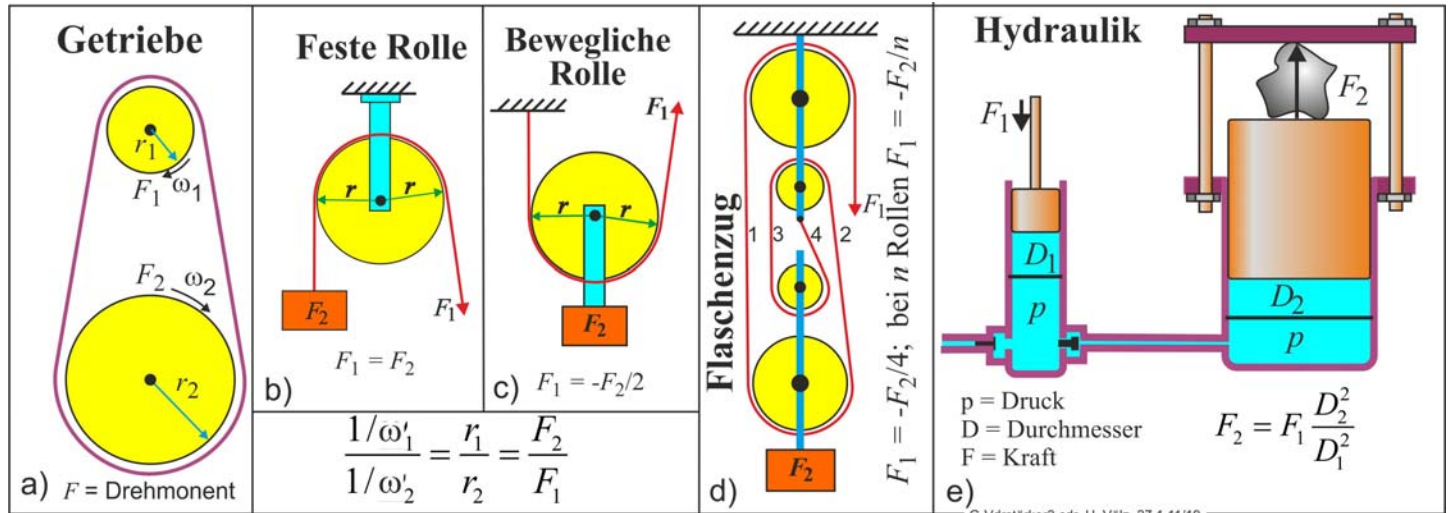
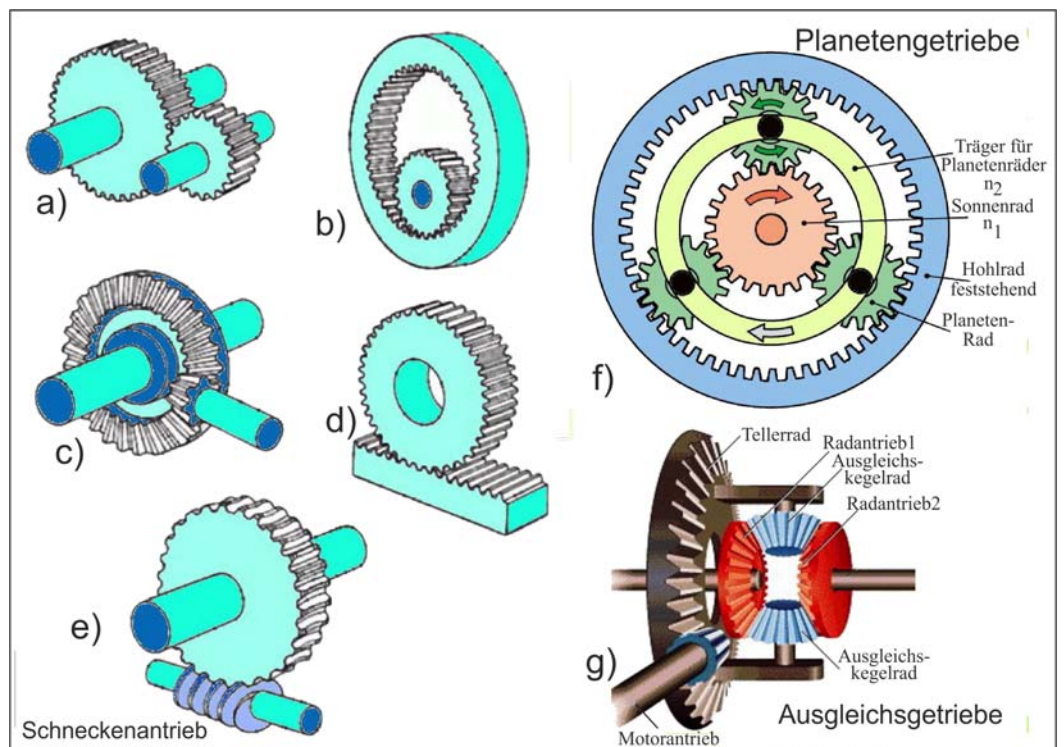


Bild 3. Zu den Anwendungen von Rollen und der Hydraulik.

Reib- und Zahnradgetriebe

Die Rollentechnik ist recht einfach und preiswert herzustellen. Außerdem kann sie auch raumsparend zerlegt und damit gut transportiert werden. Sie hat aber zwei Nachteile: 1. Kann das Seil bzw. der Riemen verschleifen und sogar zerreißen. 2. Kann auf den Rollen Schlupf auftreten. Das gilt ähnlich auch für die Reibradantriebe. Ansonsten unterscheiden sie sich im Aufbau nur relativ wenig von den Zahnradantrieben. Daher werden sie hier nur mittelbar behandelt. Es ist lediglich die Anzahl der Zähne durch den Umfang der Räder zu ersetzen. Mit beiden Abtriebarten lassen sich noch zusätzliche Ansprüche erfüllen. Je nach ihren Aufbau ermöglichen sie nicht nur Umwandlungen bezüglich Kraft und Weg, sondern auch zwischen Drehzahlen, Drehmomenten, Kräften, Drehrichtungen und Trägheitsmomenten. Einen Überblick zu den Zahnradantrieben gibt **Bild 4**. In den Beispielen (a) bis (c) werden die Drehzahlen entsprechend dem Verhältnis der Anzahl von Zähnen beider Räder geändert. Bei (b) kann das kleinere Rad auch am Außenrand der großen Rades angreifen. In allen Fällen kann entweder die Drehzahl auf Kosten der Kraft erhöht oder die Kraft auf Kosten der Drehzahl verstärkt werden. Beim Beispiel (d) erfolgt eine Umwandlung zwischen Drehung und Linearverschiebung. Der Schneckenantrieb (e) ähnelt der Schraube von Bild 2e. Sinnvoll und möglich ist er jedoch nur zur erheblichen Drehzahlreduzierung mit passender Kraftverstärkung. Ähnliches leistet auch der Planetenantrieb (f). Bei ihm zeigen jedoch die beiden Achsen in die gleiche Richtung. Das Drehzahlverhältnis ist bei ihm durch die Anzahl der Zähne vom Sonnenrad n_1 zu Planetenrädern n_2 bestimmt. Schließlich sie noch eine einfache Variante des Ausgleichsgetriebes (g) erwähnt, das auch Differentialgetriebe oder vereinfacht Differential genannt wird. In Kurvenfahrten mit vierrädrigen Fahrzeugen kann so die unterschiedliche Drehzahl der beiden hinteren nichtgelenkten Räder ausgeglichen werden.

Bild 4. Beispiele für Zahnradgetriebe.



Transformatoren

Außer dem bisher angewendeten Energieprodukt aus Kraft mal Zeit gibt in der Elektrotechnik auch Möglichkeiten aus Spannung U und Strom I gemäß $N = U \cdot I$, dabei werden vor allem Transformatoren angewendet. So kann der Strom I zu Ungunsten der Spannung U vergrößert werden bzw. umgekehrt und es sehr große Ströme oder Spannungen zu gewinnen. Dabei ist der Spartransformator mit einer Anzapfung der einen Spule dem einseitigen Hebel äquivalent **Bild 5a** und c bzw. Bild 2a. An der Anzapfung bei n_2 Windungen beträgt die Spannung nur $U_2 = U_1 \cdot n_2 / n_{\text{total}}$. Dafür gilt für die Ströme $I_2 = I_1 \cdot n_{\text{total}} / n_2$. Gebräuchlicher ist der Trenntrafo (b), bei dem zwei Wicklungen mit unterschiedlicher Drahtstärke hergestellt werden.

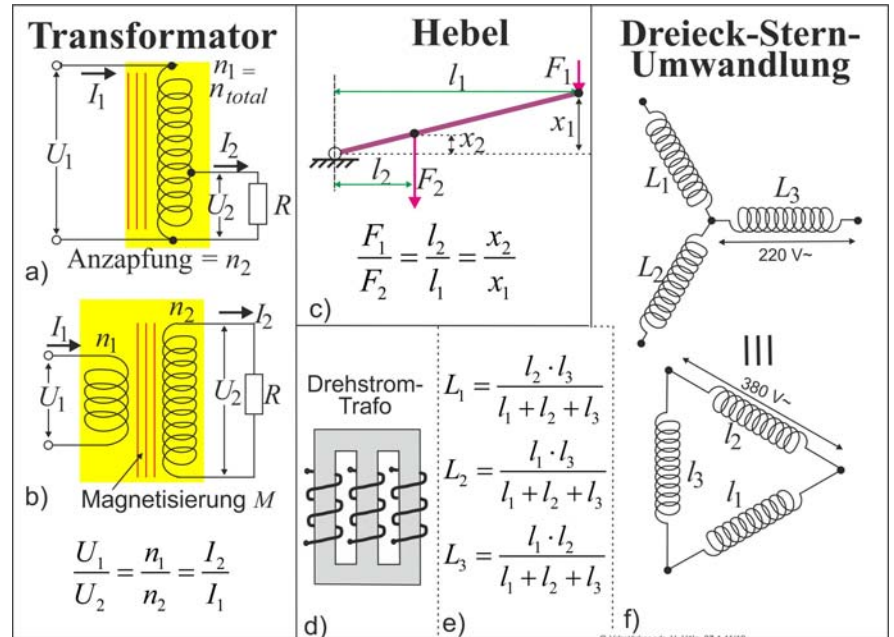
Eine andere Anwendung ist der Drehstromtrafo mit drei getrennten Wicklungen auf unterschiedlichen Kernteilen (d). Sie können in Stern- oder Dreieckschaltung betrieben werden (f). Durch die Phasenwinkel von jeweils 120° können dann 3-mal 380 bzw. 220 V erhalten werden. Der große Vorteil dieser Technik besteht hauptsächlich darin, dass so für die Fernübertragung nur drei Drähte benötigt werden. Der mittlere Nullleiter überträgt im üblichen Betriebsfall keinen Strom. Ferner lassen sich mühelos Drehfelder für Motoren mit zusätzlich umschaltbarer Leistungsaufnahme (Start- und Dauerbetrieb; Dreieck-Stern) erzeugen. Die Umrechnung zwischen den beiden Schaltungsarten erfolgt gemäß (e). Für die umgekehrte Auflösung gilt:

$$l_1 = \frac{L_1 \cdot L_2 + L_2 \cdot L_3 + L_1 \cdot L_3}{L_1}$$

$$l_2 = \frac{L_1 \cdot L_2 + L_2 \cdot L_3 + L_1 \cdot L_3}{L_2}$$

$$l_3 = \frac{L_1 \cdot L_2 + L_2 \cdot L_3 + L_1 \cdot L_3}{L_3}$$

Bild 5. Anwendungsbeispiele für Transformatoren.



3. I-Verstärker mittels Impulsen

Bei den bisherigen Beispielen wurde die Verstärkung durch die Hervorhebung eines Parameters zu Lasten eines anderen erreicht. Vor allen besteht dabei die Möglichkeit es rein mechanisch durch Integration über die Zeit erreichen. Die Kraft eines Impulses (umgangssprachlich Wucht oder Schwung) leitet sich dabei aus der Masse m und deren Geschwindigkeit v ab:

$$p = m \cdot v.$$

Die hierzu gehörende kinetische (Bewegungs-) **Energie** beträgt (in N·s = kg·m/s

$$E_{\text{kin}} = \frac{1}{2} \cdot m \cdot v^2.$$

Beim freien Fall einer Masse m im Erdfeld mit der Erdbeschleunigung $g \approx 9,81 \text{ m/s}^2$ entsteht bei einer Fallhöhe h die Geschwindigkeit

$$v = \sqrt{2 \cdot g \cdot h}.$$

Dazu gehören der Impuls

$$p = m \cdot \sqrt{2 \cdot g \cdot h} \text{ und die Energie } E_{\text{kin}} = m \cdot g \cdot h.$$

Trifft nun die Masse nun auf einen festen, harten Gegenstand, so wird dadurch die Geschwindigkeit momentan auf $v \rightarrow 0$ abgebremst. Dieser Kraftstoß überträgt dann die Gesamtenergie auf den Gegenstand, wobei eine beachtliche Impulsverstärkung auftritt. Genau das wird bei jeder Ramme ähnlich wie in **Bild 6** genutzt. Sie wurde erst relativ spät erfunden und erstmalig 1475 erwähnt. Ursprünglich dürfte sie in Kriegen zur Zerstörung von Festungen eingesetzt gewesen sein. Heute gibt es verschiedene Rammen, die sich hauptsächlich durch den „Antrieb“ unterscheiden, u. a. Schlag-, Dampf-, Explosions-, Pressluft-, Diesel- und Hydraulik-Rammen. Bei einer Vibrations-Ramme werden die Schläge periodisch wiederholt.

Während die Ramme die Schwerkraft nutzt, wird beim Hammer mit ähnlichem Ziel die manuelle Bewegung aufaddiert. Die Wirkung des **Hammers** dürfte schon in der frühen Steinzeit genutzt worden sein. Mit ihm kann mittels einer Schnur schließlich auch ein Nagel herausgezogen werden. Bei der Axt kommt durch die scharfe Schneide neben der Wirkung des Hammers noch die des Keils gemäß Bild 2d hinzu. Von der Steinaxt existieren Funde aus mindesten **6000 v. Chr.** vor.

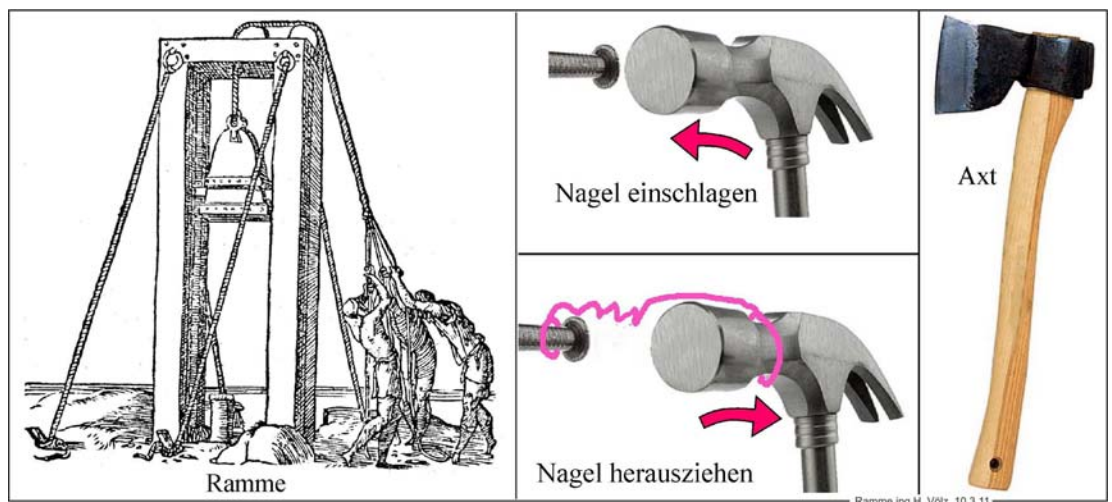


Bild 6. Impulswirkung bei Ramme, Hammer und Axt.

3. D-Verstärker beginnen mit dem Relais

D steht hier für **digital** und **diskret**. Relais stammt von französisch *relai* zurücklassen und betraf zunächst den Pferdewechsel beim Postverkehr und Militär im Sinne von Umspann-Ort, es betraf auch die an bestimmten Orten aufgestellte Reiter, die der Überbringung von Befehlen, Nachrichten dienten und schließlich den Weg zwischen Wall und Graben einer Festung. Heute wird damit Bauelemente bezeichnet, die als elektrisch betriebene Schalter sind. Formal weist ihre Verwendung als Verstärker **Bild 7b** aus. Darin ist aber nur die zum Schalten geringere Kraft ausgewiesen, aber auch der manuell geschaltete Strom ist erheblich geringer.

Bild 7. Zur Veranschaulichung des Relais als Verstärker.

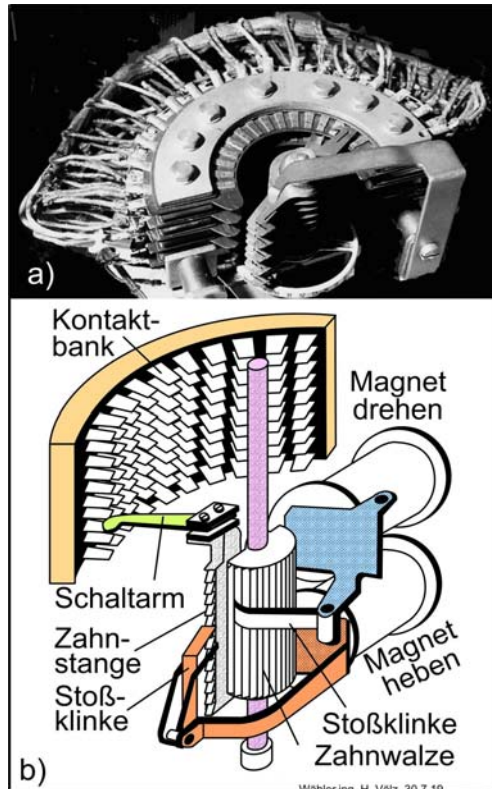
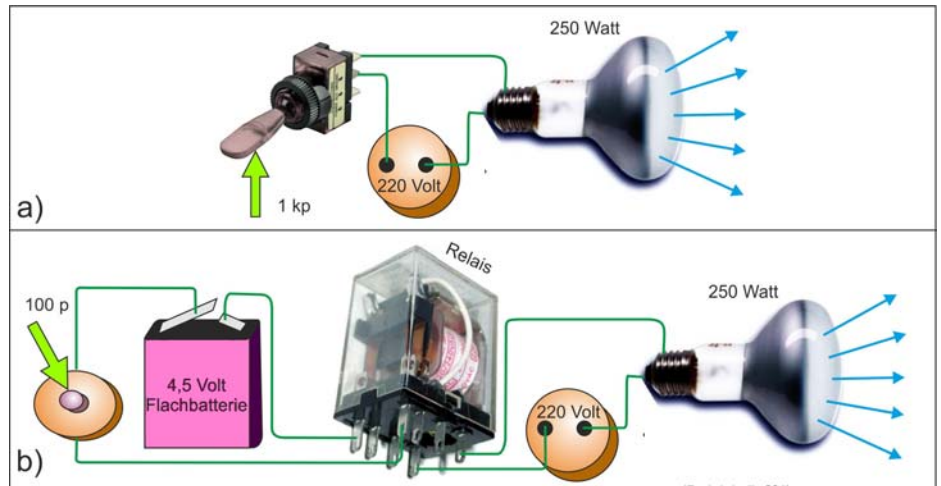


Bild 8. a) Drehwähler, b) Heb-Dreh-Wähler.

1837	Samuel Finley Breese MORSE baut das erste Relais
1838	Charles Wheatstone baut Relais für Telegrafie
1883	Industrielle Anwendungen von Relais
1908	Erstes Relais der Fernsprechtechnik , ähnlich Bild 8b
1925	Thyatron
1927	Das Flachrelais wird eingeführt
1936	Relais-Rechner Z1 von Konrad Zuse in der elterlichen Wohnung fertig gestellt
1936	Reed-Kontakt von W. B. Elwood (Bell Laboratories) patentiert, wird aber erst ab 1950 produziert (Bild 9d)
1957	Thyristor

Ein **Relais** hat bei seiner Anwendung **Grenzen**. Es kann nicht sehr schnell Schalten, die Kontakte können korrodieren und bei größeren Schaltleistungen können sie durch Funkenbildung unbrauchbar werden. Als Verbesserungen entstanden das **Thyatron** und die verschiedenen Varianten der **Thyristoren**.

Das **Thyatron**¹ geht auf die Gasentladung – u. a. bei Glimmlampen und Leuchtstoffröhren – zurück. Ein Rohr ist mit Quecksilberdampf, Xenon, Neon, Krypton oder Wasserstoff gefüllt (10a). Das typische Verhalten hängt von der Gasart und dem Gasdruck ab (b). Die Strom-Spannungskennlinie ist recht komplex (d). Es gibt 3 Entladungsgebiete: die fast unsichtbare Dunkelentladung, das schwache Glimmlicht und die helle, stromstarke Bogenentladung (d). Mit einer geheizten Kathode entstanden daraus die Quecksilbergleichrichter, die in einer Stromrichtung zünden, eine Halbwelle sperren. Um 1925 entstand daraus analog der Triode (s. u.) das Thyatron. Mit dem Gitter wird es I_{st} gezündet (Kreise mit Pfeil in e). Beim Wechsel der Stromrichtung bricht die Bogenentladung. So wurde ähnlich dem Relais ein Ein- und Ausschalten möglich. Es wurde bis zu beachtli-



Die „Verstärkung“ ermöglichte die erste Anwendung des Relais in der Telegrafieübertragung auf sehr langen Leitungen. Weitaus umfangreicher erfolgte die Anwendung in der Fernsprechtechnik für die automatische Wahl von Telefonnummern. Dafür entstanden vor der diskreten Dreh- und später der Hebdrehwähler (**Strangwar**) (**Bild 8**).

Sehr schnell wurde dann das „ursprüngliche“ Relais äußerst vielfältig eingesetzt und dafür ständig weiterentwickelt und auch erheblich verkleinert. So entstanden u.a. das typische Rund- und Flachrelais (**Bild 9a, b**), hauptsächlich existieren die folgenden Betriebsvarianten:

1. Nur bei Strom-Erregung einschaltend (Bild 9a, b).
2. Durch Strom-Impulse erfolgt Umschalten, z. B. Telefon-Relais (Bild 8a, b).
3. Mittels Strom erfolgt nur eine Umschalten für 2 Zustände (Telegraphenrelais nach Bild 9c).
4. Aus- oder Umschalten durch äußere Magnetfelder z. B. Reed-Kontakte (Bild 9d)
5. Relais mit einem oder mehreren Schaltkontakten (Bild 9a, b)
6. Bei großen elektrischen Leistungen werden Schütze benutzt. Sie enthalten zusätzliche eine Funkenlöschung.
7. Elektronische Relais, u. a. Thyratrons, Thyristoren, S. Kapitel ###.

Eine spezielle aber hoch leistungsfähige Anwendung erfuhren die Relais bei den sehr frühen Rechnern. Insgesamt war so die Relaisstechnik ganz wesentlich für den Übergang zur Elektronik mit den Elektronenröhren und Transistoren.

Die wichtigsten Geschichtsdaten, einschließlich der dazu gehörenden elektronischen Varianten, fasst die folgende Tabelle zusammen.

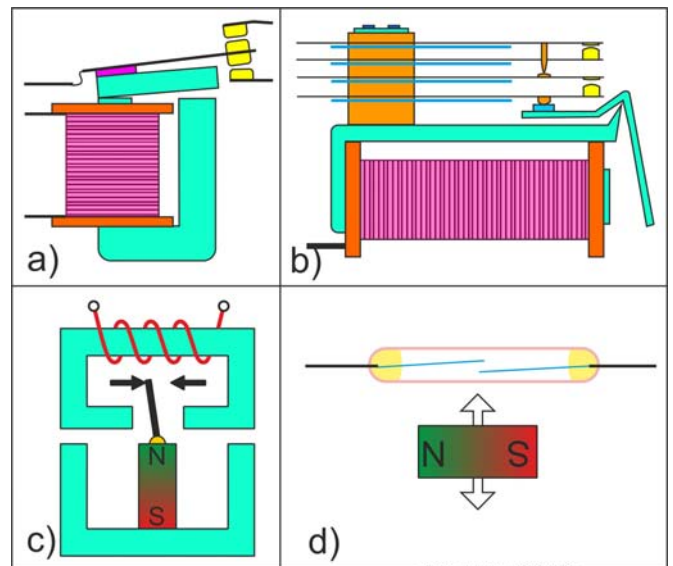
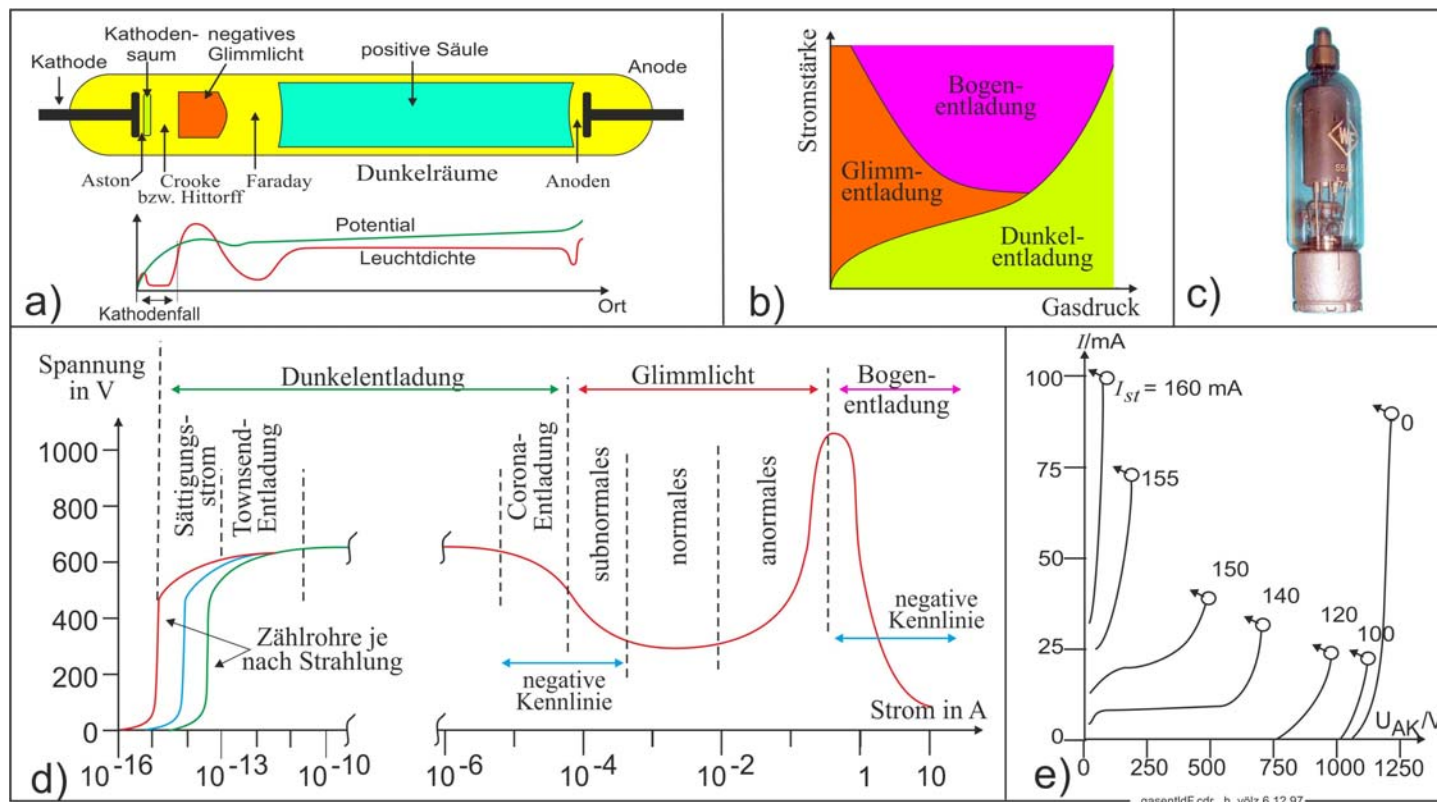


Bild 9. Beispiele für klassische Relais.

¹ Griechisch Thyra Tor, großer Stein als Tür vor den Eingang gelegt, großer türförmiger Schild.

chen Leistungen möglich, z. B. für 5kV/6A vom Werk für Fernsehelektronik (DDR) (c). Teilweise wurden solche Röhren auch Ignotron, Ionenröhre (ion tube) oder Stromtor genannt. Sie wurden ab etwa 1957 vollständig durch ähnlich wirkende Halbleiterbauelemente abgelöst.



10. Von der Gasentladung zum Thyatron,

Der **Thyristor** beruht auf Halbleiter (s. u.). Das Wort vereint Thyristor und Transistor. Gemäß Bild 10 existiert er in vielen Varianten, die sich aus mehreren und „nebeneinander gelegten“ pn-Übergängen ableiten lassen. Dabei werden Dioden (a) und Trioden (b) sowie im in „rückwärtigen“ Zweig drei Kennlinien unterschieden (c und d). So ergeben sich u. a. **TRIAC**, **IGBT**, **IGCT** (integrated gate x Thyristor), GTO-Thyristor (Gate Turn Off) und SCR. Eine wichtige Anwendung ist die Phasenanschnittsteuerung gemäß dem Schaltbild (h). Dabei werden bidirektionale Thyristoren verwendet. So kann die Helligkeit von Glühlampen beliebig gesenkt (gedimmt) werden. Dabei werden mit dem den Steuerimpuls durch Verstellen des Reglers nur die grün gezeichneten Abschnitte der Sinuskurve durch gelassen (e bis f). Ein typisches Bauelement zeigt (i). Üblicherweise ist das Dimmen nicht für Leuchtstoffröhren und LED möglich.

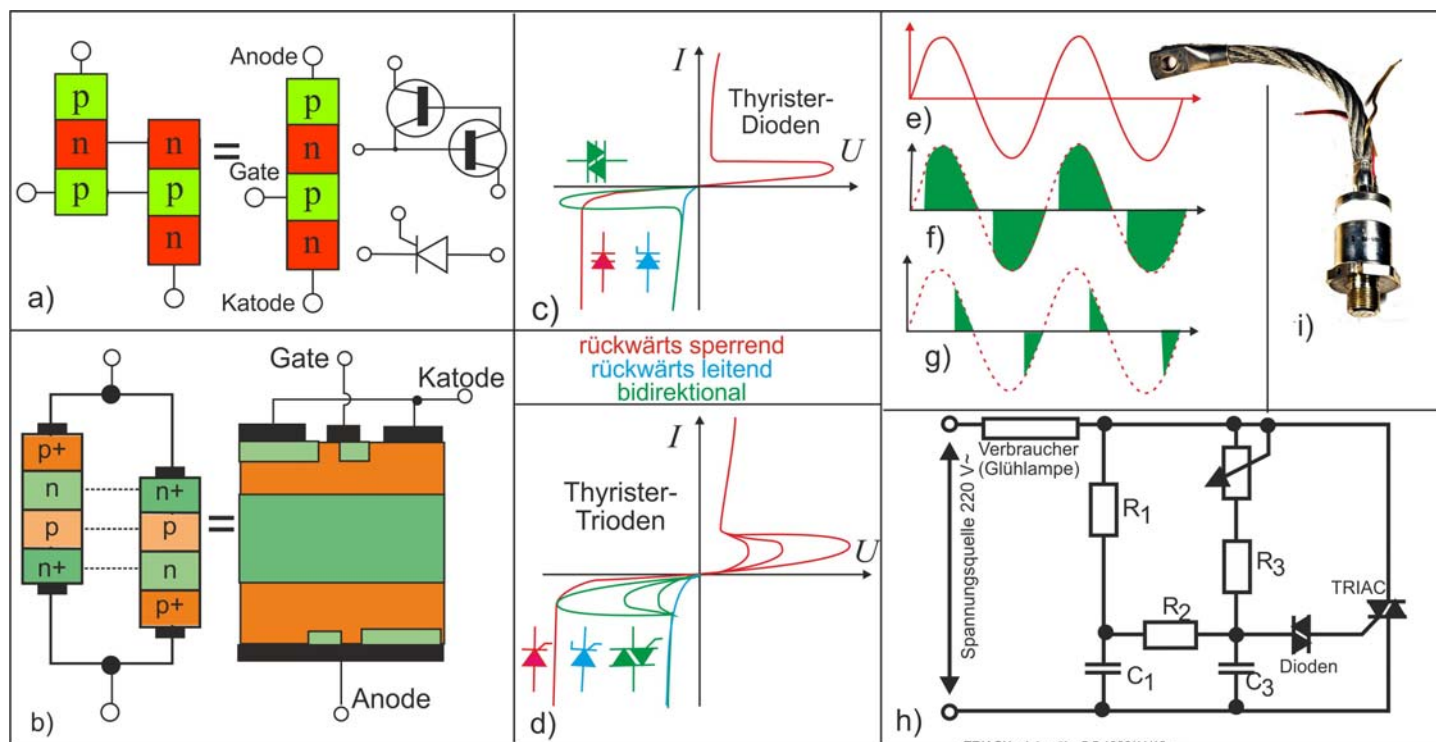
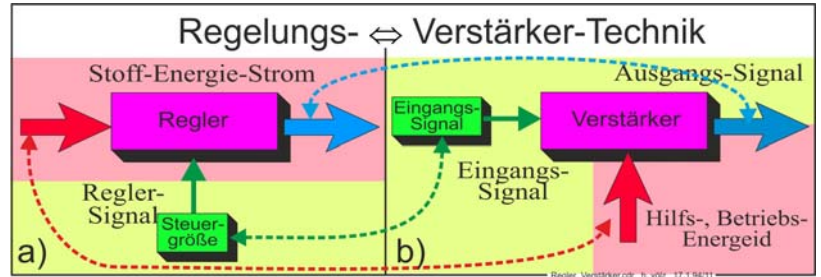


Bild 10. Varianten der Thyristoren.

4. S-Verstärker für elektrische Signale

S-Verstärker sind fundamental für die gesamte Elektronik. Die Signale, die sie mittelbar verstärken, sind Strom - oder Spannungsänderungen. Dabei benötigen sie eine Hilfs-, Betriebsenergie die verändert wird. Das weisen die beiden Systemdarstellungen von **Bild 11** deutlich aus. Formal unterscheiden sich beide nur durch Vertauschung der Anschlüsse gegenüber der Regelung (a) vor (vgl. Kapitel ###). Es gibt bisher nur 3 Bauelemente: digital die Relais (s. o.) und kontinuierlich Röhren und Transistoren

Bild 11. Vergleich von elektrischen Verstärker und kybernetischem Regler.



Elektronenröhren

Mit der Erfindung der Elektronenröhre entstand erstmalig die Möglichkeit, eine Energie fast leistungslos zu beeinflussen. Rein theoretisch entspricht das einer unendlichen (Leistungs-) Verstärkung und bedarf unbedingt einer Erklärung². In **Bild 12** wird dafür ein Analogie betrachtet (a). Auf einem Berg befindet sich ein Stausee. Über einer einstellbaren Barriere wird so eine regelbare Wassermenge zum Kraftwerk im Tal geleitet. Analoges erfolgt bei der Röhre (b). Der Heizfaden (Katode) wird durch elektrischen Strom auf eine hohe Temperatur erhitzt. Dadurch treten an seiner Oberfläche Elektronen aus und erzeugen eine Elektronenwolke, deren Geschwindigkeitsverteilung (c) zeigt. Von der positiven Anode werden die Elektronen angezogen und daher fließt der Anodenstrom. Zur Steuerung ist zwischen dem Heizfaden und der Anode ein „löchriges“ Gitter angeordnet. Wird es negativ aufgeladen, so können nur einige Elektronen die Gitterbarriere überwinden. Je nach der Gitterspannung ist das nur für die schnellen (c) möglich. Da hierbei kein Gitterstrom auftreten kann, erfolgt die Steuerung leistungslos. Obwohl die ersten Röhren seit 1904 erfolgreich benutzt wurden, erfolgte erst 1928, als die sehr erfolgreiche Anwendung offensichtlich war, die Verleihung des Nobelpreises (s. u.).

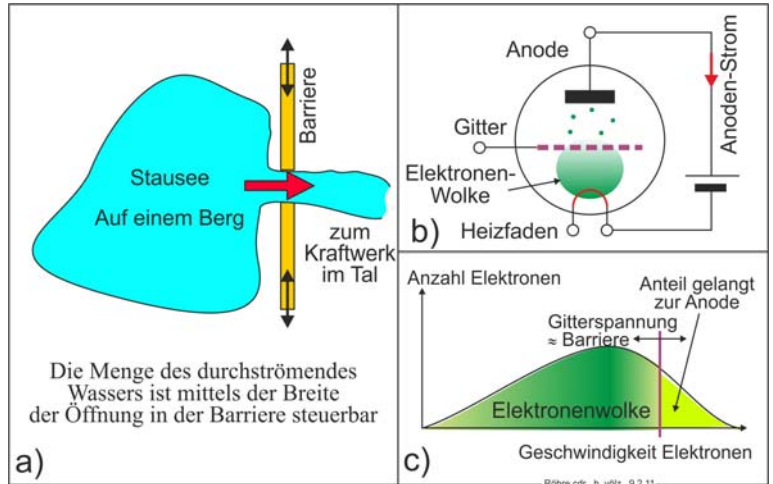


Bild 12. Zur Erklärung der leistungslosen Steuerung des Anodenstromes entsprechend der negativen Gitterspannung.

Für die Anwendung der Röhren – sie wird wegen der 3 Abschlüsse Triode genannt – sind und das Schaltbild ihre Kennlinien wesentlich (**Bild 13**). Die Eingangskennlinie beschreibt den Zusammenhang von Gitterspannung U_g und den Anodenstrom I_a mit der Anodenspannung als U_a als Parameter. Die Ausgangskennlinie zeigt I_a als Funktion von U_a mit U_g als Parameter. In ihr lässt sich die Arbeitskennlinie für R_A ergänzen. Sie zeigt an, wie stark Anodenspannung durch den Abfall an R_A sinkt. Mit ihrer Hilfe kann dann die Schwingungsverstärkung gemäß **Bild 14** aufgezeigt werden. Daraus wird deutlich, dass der notwendige Spannungsabfall am R_A die Betriebsspannung U_b um $R_A I_a$ reduziert und damit zurückwirkend die Verstärkung V gegenüber den rein theoretischen Wert absenkt.

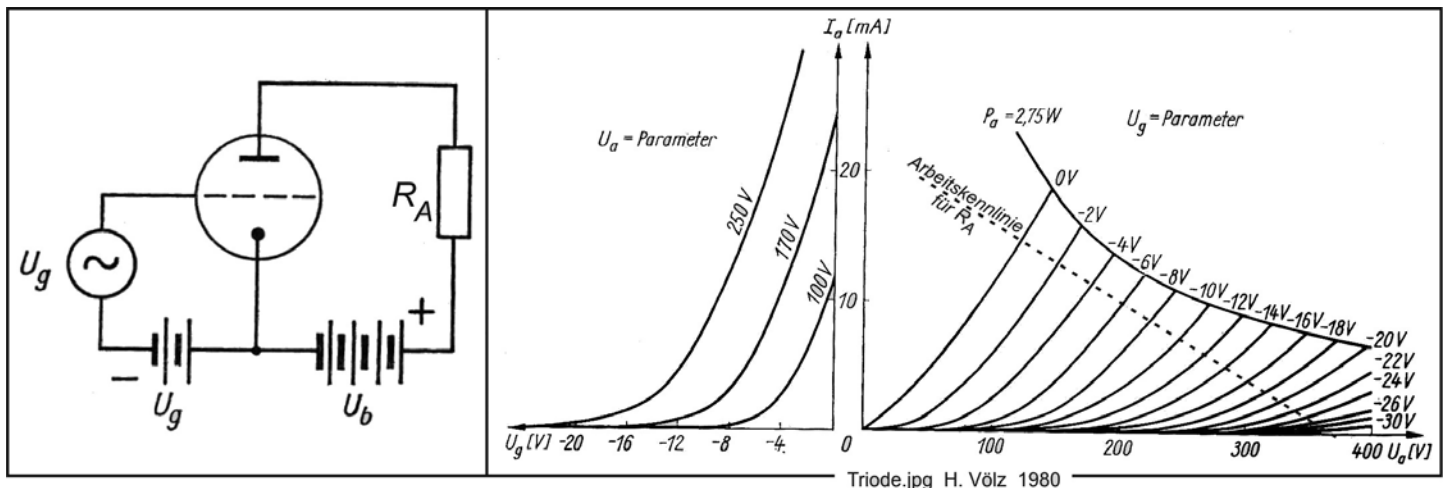


Bild 13. Schaltung (links) und typische Kennlinien einer Triode.

Die Penthode

Der Nachteil der Rückwirkung der Ausgangsspannung U_a auf die wirklich anliegende Anodenspannung U_b führte zur Einführung eines zweiten Gitters, des Schirmgitters G_2 , das mit einer unveränderlichen Spannung U_{G2} versorgt wird (**Bild 15**). So kann der veränderliche Anodenstrom kaum noch auf zurückwirken. Um auch Sekundärelektronen von der Anode unwirksam zu machen, wurde schließlich noch das Bremsgitter G_3 eingefügt und auf Nullpotential gelegt. So entsteht der Potentialverlauf in der Röhre nach (c). Die grünern Flächen zeigen die Unterschiede zwischen den Gitterdrähten an. Als Ergebnis zeigt deutlich den stark abgeflachten Verlauf bei den Ausgangskennlinien (d).

² Zu den Verstärkerröhren ist heute kaum noch Literatur verfügbar. Lediglich in der ersten Auflage meines Elektronikbuches sind die Grundlagen und Anwendungen noch sehr ausführlich enthalten [Völ74].

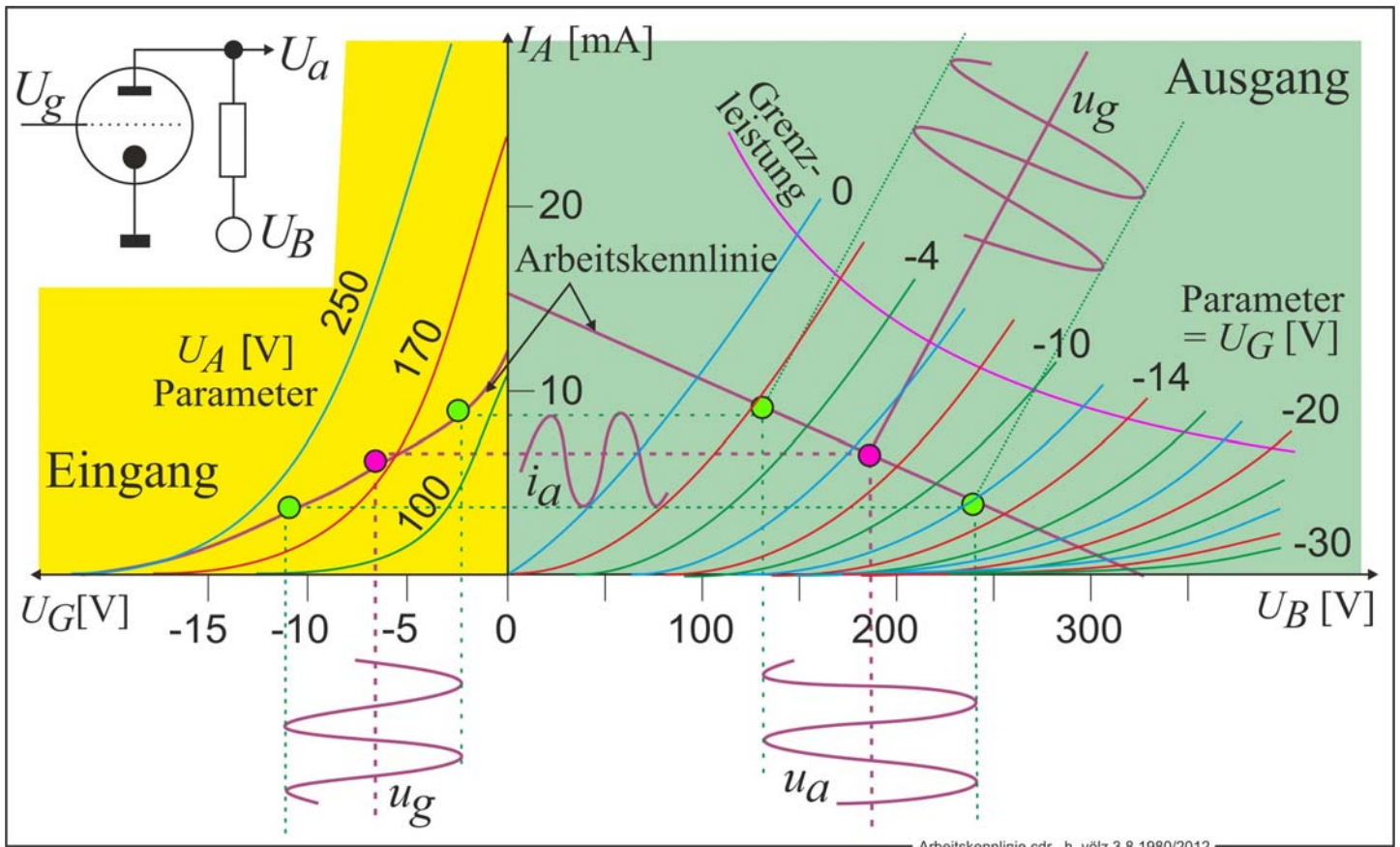


Bild 14. Schaltbild und Signale bei der Triode.

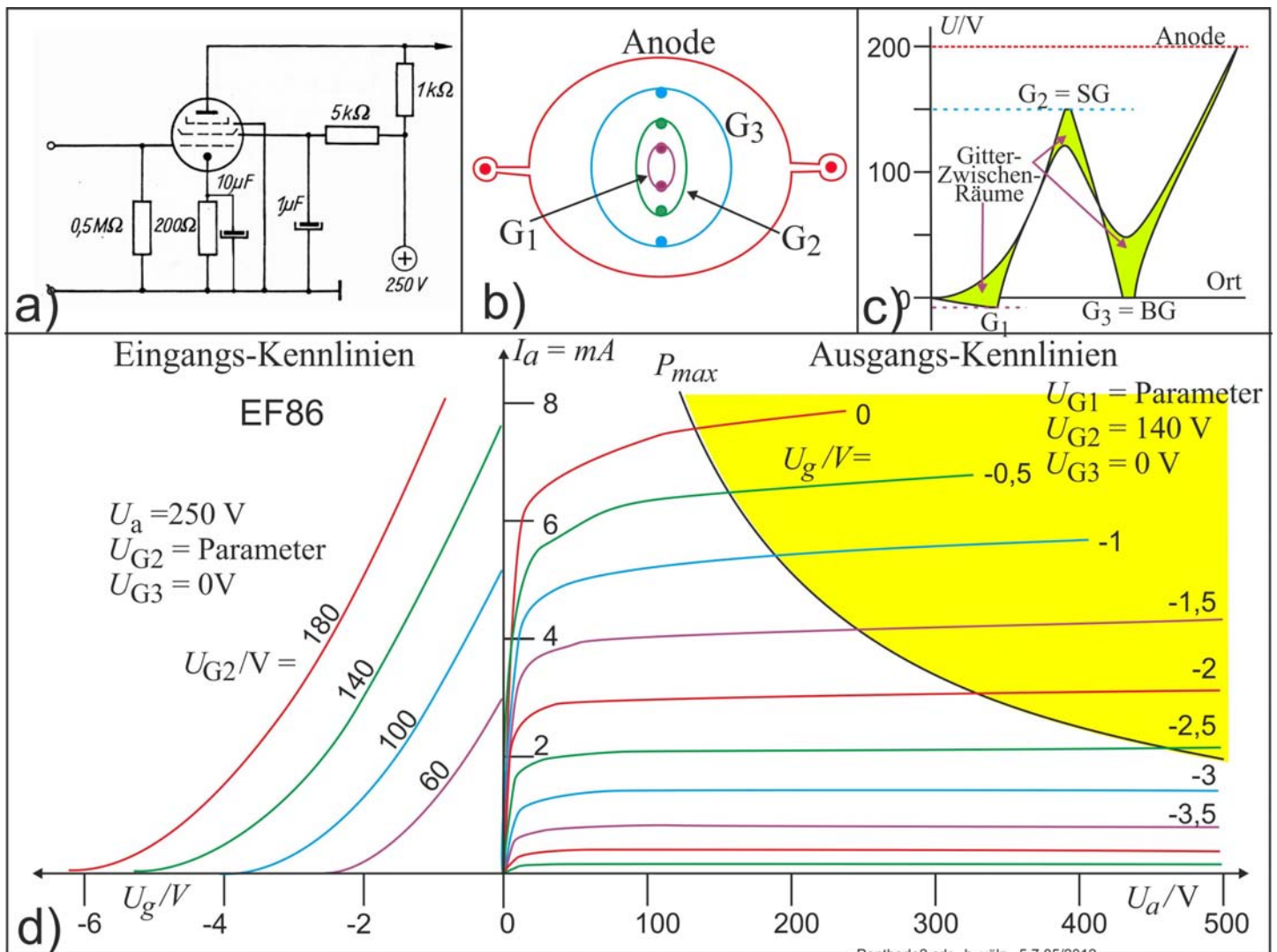
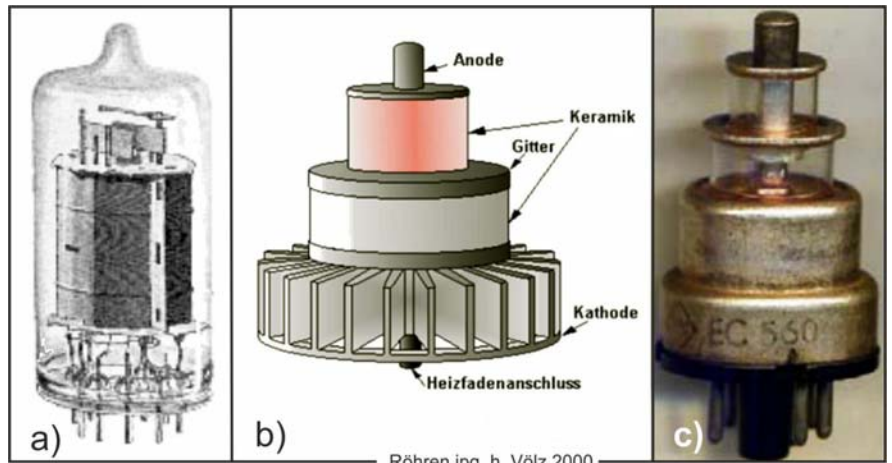


Bild 15. Schaltung, Aufbaustruktur, Feldverlauf und Kennlinien einer Penthode.

Weitere Verstärkerröhren

Es sind mehrere weitere Varianten entstanden, u. a. die Dioden für die Gleichrichtung und Demodulation, die Hexode zur Mischung (Superhet), die Endröhren mit erhöhter Leistung, Röhren mit mehreren Teilröhren, insbesondere die Dreifachröhre von Ardenne/Loewe. Dann Spezialröhren für höchste Frequenzen (s. Abschnitt ###). Nur ein Sonderaufbau sei noch kurz beschrieben. Bei hohen Frequenzen ist die Laufzeit der Elektronen von der Kathode zur Anode beachtlich. Daher ergibt sich eine Grenzverstärkung von mehreren MHz. Einen Ausweg ist die Scheibendiode, bei der durch die ebene Ausführung der Gitter-Kathoden-Abstand auf $<10\mu\text{m}$ verkleinert ist. So wurden 10 GHz möglich. Neben der Laufzeitverzögerung begrenzten auch restliche Kapazitäten und Induktivitäten die höchste Frequenz. **Bild 16b** und c zeigen eine Scheibentriode und zwar im Vergleich (nicht maßstabsgerecht) mit einer relativ späten üblichen Verstärkerröhre.

Bild 16. Nicht maßstabgetreuer Vergleich einer Scheibentriode (b, c) mit einer üblichen



Verstärkerröhre (a).

Geschichte der Röhren

1883 THOMAS ALVA EDISON (1847 - 1931) entdeckt **Glüh-effekt**, Elektronenaustritt, Zusatzelektrode (Diode)
 1879 WILLIAM CROOKES (1832 - 1919) beschreibt wesentliche Eigenschaften der **Katodenstrahlen**
 1897 KARL FERDINAND MARQUIS DE BRAUN (1850 - 1918) erfindet mit JONATHAN ZENNECK die **Katodenstrahlröhre**
 1901 OWEN WILLIAMS RICHARDSON (1879 - 1959) Gesetz der **Glühemission**
 1906 ROBERT VON LIEBEN (1878 - 1913) Patent **quecksilbergefüllte Verstärkerröhre, 2 Elektroden**, elektrische oder magnetische Beeinflussung
 1902 PETER COOPER-HEWITT patentiert **Quecksilberdampfgleichrichter** (nicht zur Verstärkung geeignet)
 1904 JOHN AMBROSE FLEMING (1849 - 1945) patentiert **Vakuum-Diode**
 1906 LEE DE FOREST (1873 - 1961) meldet **gasgefüllte Audionröhre** mit zusätzlichem Gitter zum Patent an, hat kaum Bedeutung erlangt
 1912 FOREST von Bell Telephone Laboratories stellt **Röhrenverstärker** vor
 1913 Röhrenverstärker für **Telefonverbindungen** zwischen New York und Baltimore
 1914 Röhrenverstärker für **Atlantik-Seekabel**
 1916 Bei Siemens & Halske entwickelt WALTER SCHOTTKY (1886 - 1976) **Tetrode** (Pentode jedoch ohne Bremsgitter)
 1924 Französische Firma Métal produziert **Doppelgitter-Röhre** als Mischröhre für Radioempfänger (Vorstufe der Hexode)
 1926 BERNARD D. H. TELLEGEN (1900 - 1990) entwickelt bei Philips die **Penthode** zur Serienreife
 1926 MANFRED BARON VON ARDENNE (1907 - 1997) mit SIEGMUND LOEWE Mehrsystemröhren = **Dreifachröhre**
 1928 RICHARDSON **Nobelpreis** für Physik (Glühemission)

Transistoren

Einen Transistor gibt es seit 1947, also 22 Jahre nach der Elektronenröhre. Seine Grundlagen sind die Halbleiteereigenschaften und deren Grundlage ist das Bändermodell. Es folgt wiederum aus den Eigenschaften des Atomaufbaus nach der Quantenphysik (**Bild 17**). In der klassischen Beschreibung bewegen sich die Elektronen eines Atoms auf unterschiedlichen Bahnen. Durch quantentheoretische Einschränkungen kann das nur auf ausgewählten Energieniveaus (Bahnen) der Schalen K, L, M usw. geschehen. Die Schalen werden von unten mit Elektronen aufgefüllt. In der n -ten Schale können sich dabei maximal $2n^2$ Elektronen befinden. Wegen des Pauli-Verbots müssen sie sich aber quantenphysikalisch unterscheiden. Das wird zusätzlich durch ihre Nebenquantenzahlen s, p, d usw. (klassisch: Kreis- und verschiedene Ellipsenbahnen) sowie durch ihren Spin ($\pm \frac{1}{2}$) erfasst. Dieser Zusammenhang wird durch einen **Potentialtopf** (a) beschrieben. Erreichen Elektronen durch externe Anregung eine hinreichend große Energie, so können sie das Atom verlassen, und es tritt Ionisierung ein.

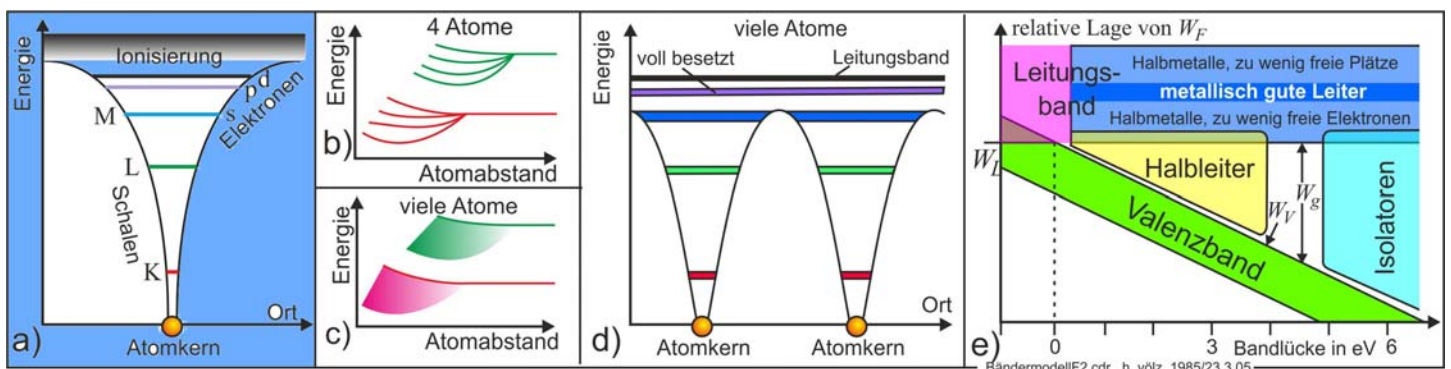


Bild 17. Zu den quantenphysikalischen Grundlagen der Halbleiter.

Werden mehrere hinreichend Atome dicht – z. B. in einem Kristall – angeordnet, so tritt eine erhebliche Wechselwirkung auf, und sie werden ein einziges Quantensystem. Dann müssen Elektronen mit gleicher Haupt- und Nebenquantenzahl unterschiedliche Energien (u. a. durch Wärme oder Bewegung) annehmen. So spalten sich die Bahnen auf. Für vier Atome zeigt das bezüglich des gegenseitigen Atomabstandes (b). Sind viele Atome vorhanden, so werden aus den singulären Bahnen Bänder (c). Für zwei ausgewählte benachbarte Atome zeigt das (d) durch „verknüpfte“ Potentialtöpfe. U. a. entstehen so das Leitungs- und das Valenzband (e). Im **Valenzband**³ befindet sich die letzte voll besetzte Schale. Es reicht bis zur maximalen Energie W_V und bestimmt im Wesentlichen die chemischen Eigenschaften der Elemente. Das **Leitungsband** ist das nächsthöhere Band. Es enthält keine Elektronen oder ist nur unvollständig mit Elektronen gefüllt. Je nachdem sind mehr oder wenige gute Leiter. Seine untere Energie-

³ **Valenzband** von lateinisch valentia Stärke, Kraft, und von valere stark, gesund sein.

grenze ist W_L . Die in ihm vorhandenen Elektronen gehören zum gesamten Kristall. Das wird klassisch durch das Drudensche Elektronengas beschrieben. Zwischen beiden Bändern besteht meist eine **Bandlücke** (engl. gap) mit dem Energieabstand W_g . Sie ist durch Kristallmaterial bestimmt. Bei sehr geringem Bandabstand liegt ein Leiter vor. Für Bandlücken von etwa 0,5 bis 3,5 eV⁴ existieren die Halbleiter. Bei Bandlücken >5 eV treten Isolatoren auf. Außerdem ändert sich mit dem Bandabstand erheblich die Temperaturabhängigkeit der Leitfähigkeit. Bei den Leitern ändert sie sich nur relativ wenig; bei Halbleitern dagegen über viele Zehnerpotenzen etwa proportional zu e^T und bildet so eine Gerade bezüglich $1/T$.

Bei einem fehlerfreien Monokristall wirkt die thermische Energie statistisch auf die Elektronen des Valenzbandes. Erreicht dabei ein Elektron eine Energie größer als der Bandabstand W_g , so gelangt es ins Leitungsband. Das wird (Träger-) **Generation** genannt (**Bild 18a**). Als freier Ladungsträger erhöht es dadurch die Leitfähigkeit des Materials. Bei Leitern steigt sie jedoch nur wenig mit der Temperatur an. Es können auch Elektronen des Leitungsbandes wieder zurück ins Valenzband „fallen“, dann liegt eine **Rekombination** vor.

Wird der Kristall mit Fremdmaterial legiert (dotiert), so entstehen in der Bandlücke zusätzliche Terme. Liegen sie in der Nähe des Leitungsbandes so wird von **Donatoren**⁵ gesprochen (b). Im anderen Fall sind es **Akzeptoren**⁶ (c). Die Elektronen der Donatoren spenden infolge des geringen Abstandes zusätzliche Elektronen ins Leitungsband und erhöhen die Leitfähigkeit des Materials. Die leeren Terme infolge der Akzeptoren entziehen dagegen den Valenzband Elektronen. Dadurch entstehen leere Terme, die auch Löcher genannt werden und so virtuellen positiven Elektronen entsprechen. Ähnlich den Elektronen können sie im Material ihren Ort wechseln. Das wird Löcherleitung genannt. Je nachdem, ob Elektronen beim Leitungsband oder Löcher beim Valenzband überwiegen, wird von Elektronen- oder Löcherleitung bzw. **n-Leitung** oder **p-Leitung** gesprochen.

Freie Terme können auch durch **Kristallfehler** auftreten. Da sie technologisch kaum beherrschbar sind, muss bei der Herstellung von Halbleitern von fehlerfreien Kristallen ausgegangen werden.

Die Dotierungen bewirken eine Störstellenleitung und damit deutlich veränderte **Temperaturabhängigkeiten**. Den Einfluss von Donatoren zeigt Bild 18d. Durch günstige gewählte Dotierung kann ein Bereich ΔT erzeugt werden, in dem die Ladungsträgerzahl n_e , also die Leitfähigkeit nahezu konstant bleibt. Hier werden alle Akzeptoren- bzw. Donatoren-Niveaus aufgebraucht.

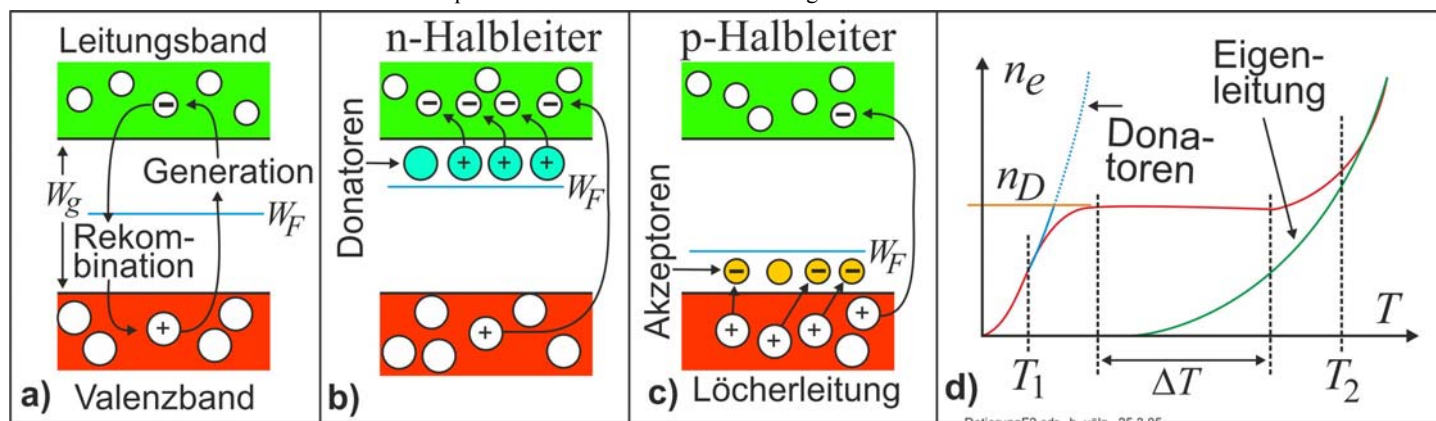


Bild 18. Zur Dotierung der Halbleiter. Die offenen Kreise im Valenz- und Leitungsband bedeuten unbesetzte Terme als mögliche Zustände für Elektronen. W_F ist betrifft die Energie des mittleren neutralen Niveaus, Fermi-Niveau.

Bei den Leitungsvorgängen sind folgende Varianten zu unterscheiden:

- **Eigenleitung** (intrinsisch) bei einem nicht dotierten Material.
- **Störstellenleitung** (extrinsisch). Sie wird durch Donatoren und Akzeptoren bewirkt.
- **n-Leitung** durch freie Elektronen im Leitungsband
- **p-Leitung** durch Löcher im Valenzband.
- **Majoritätsleitung entsprechend der** überwiegenden Anzahl vorhandener Ladungsträger (Elektronen- oder Löcherleitung).
- **Minoritätsleitung**, bestimmt durch die in kleinerer Zahl der vorhandenen Ladungsträger.

pn-Übergang

Mittels Dotierung ist entweder p- oder n-leitendes Halbleitermaterial herzustellen. Eine neue Qualität entsteht, wenn beide Materialien miteinander in Kontakt kommen, insbesondere in einem gemeinsamen Kristall. Zunächst sei angenommen, dass sie sich in geringem Abstand gegenüberstehen (**Bild 19a**). Sobald der Abstand extrem klein wird, muss es ein gemeinsames Fermi-Niveau W_F geben (b). Als Folge der unmittelbaren Nachbarschaft werden dann die benachbarten Teile der Leitungs- und Valenzbänder beider Materialien gegeneinander verschoben (d). Für die Trägerdichten der Störstellen- und Eigenleitung ergibt sich der Verlauf (c). In der Umgebung der Berührungszone des eigentlichen pn-Übergangs diffundieren⁷ Elektronen von rechts nach links und Löcher entgegengesetzt. Das ist jedoch nur begrenzt möglich. Denn dabei baut sich in der Umgebung des pn-Kontaktes eine Raumladung auf (d, e). Typisch ist hierfür eine Länge von 0,1 μm . Die Raumladung wirkt dann der Diffusion entgegen und es entsteht ein dynamisches Gleichgewicht mit einem Feldstärke- und Potentialverlauf und dem Diffusionspotential

$$U_D = \frac{k \cdot T}{e_0} \ln \frac{n_i^p}{n_i^n}.$$

Es liegt bei etwa 0,2 V. In der Gleichung bedeuten k Die Boltzmann-Konstante, T die absolute Temperatur, e_0 die Elektronenladung und n_i die Trägerdichten. Leider ist U_D nicht experimentell messbar. Durch die dafür extern anzulegenden Leiter treten nämlich wesentlich größere Kontaktpotentiale auf. Das gilt sogar, obwohl der typische Diffusions- als auch Feldstrom bei einigen kA/cm² liegt.

Den typischen Aufbau eines pn-Übergangs der Halbleiterdiode zeigt (f). Ein n-dotierter Si-Einkristall wird von oben durch Diffusion mit einem Akzeptor p-dotiert. Es entsteht eine dünne p^+ -Schicht ($^+$ bedeutet hoch dotiert). Um störende Kanteneffekte zu vermeiden, wird anschließend seitlich rundherum ein Teil dieser Schicht weggeätzt. Dann erfolgen die beiden Al-Kontaktierungen an den Enden. Unter dieser Struktur sind von links nach rechts eine Spitzendiode, ein pn-Übergang, eine Röhrendiode und das Schaltbild in gleicher Polung gezeigt.

Für die Anwendung der pn-Diode ist die Strom-Spannungskennlinie (g) wichtig. Eine angelegte Spannung U wird fast nur innerhalb der Raumladungszone wirksam. Infolge der dort geringen Ladungsträgerdichte existiert dort der deutlich überwiegende Widerstand. Dabei setzt sich die wirksame Spannung – unter Beachtung der Polarität als $U_D \pm |U|$ zusammen. Dabei werden die Feldstärke, Raumladung, Ausdehnung der

⁴ 1 eV (Elektronen-Volt) $\approx 1,6 \cdot 10^{-19}$ Joule.

⁵ **Donator** lateinisch *dotare* von *dos* Mitgift, Gabe, *dare* geben, reichen, spenden und **donator** Geber, Spender, Stifter.

⁶ **Akzeptor** lateinisch *acceper* Empfänger.

⁷ **Diffundieren** lateinisch *dis* auseinander, zwischen und *fusum* gießen, fließen lassen.

Raumladezone und Verbiegung der Bänder verändert. Konstant bleibt lediglich der feldunabhängige Diffusionsstrom. Für den extern messbaren Strom gilt

$$I = I_{Sp} \left(e^{\frac{q_0 U}{kT}} - 1 \right). \quad (1)$$

Dies ist die grundlegende Gleichung der Halbleitertechnik mit dem Verlauf (g). I_{Sp} ist ein Sättigungssperrstrom, der von Eigenschaften des pn -Übergangs abhängig ist. Für kleine Si-Dioden beträgt er bei Zimmertemperatur etwa $10^{-3} \mu A$. Die Größe kT/e_0 beträgt bei Zimmertemperatur etwa 26 mV. Je nach der Polung der angelegten Spannung werden die Durchlass- und Sperrrichtung (rechts bzw. links) unterschieden. In (g) ist noch die Schwellenspannung U_S eingezeichnet. Für vereinfachte Betrachtungen die Kennlinie damit aus zu zwei Geraden und einen Ersatzwiderstand R_{Er} genähert werden:

$$I = 0 \text{ für } U < U_S \quad \text{und} \quad I = (U - U_S)/R_{Er} \text{ für } U \geq U_S.$$

Abweichende Kennlinien – vor allem bezüglich der Schwellenspannung – besitzen Dioden aus anderem Material z. B. Ge und GaAs. Bei den **Schottky-Dioden** wird der pn -Übergang durch einen Metall-Halbleiterkontakt realisiert. Bei **PIN-Dioden** ist der pn -Übergang durch ein eingefügtes Eigenleitungsgebiet verlängert

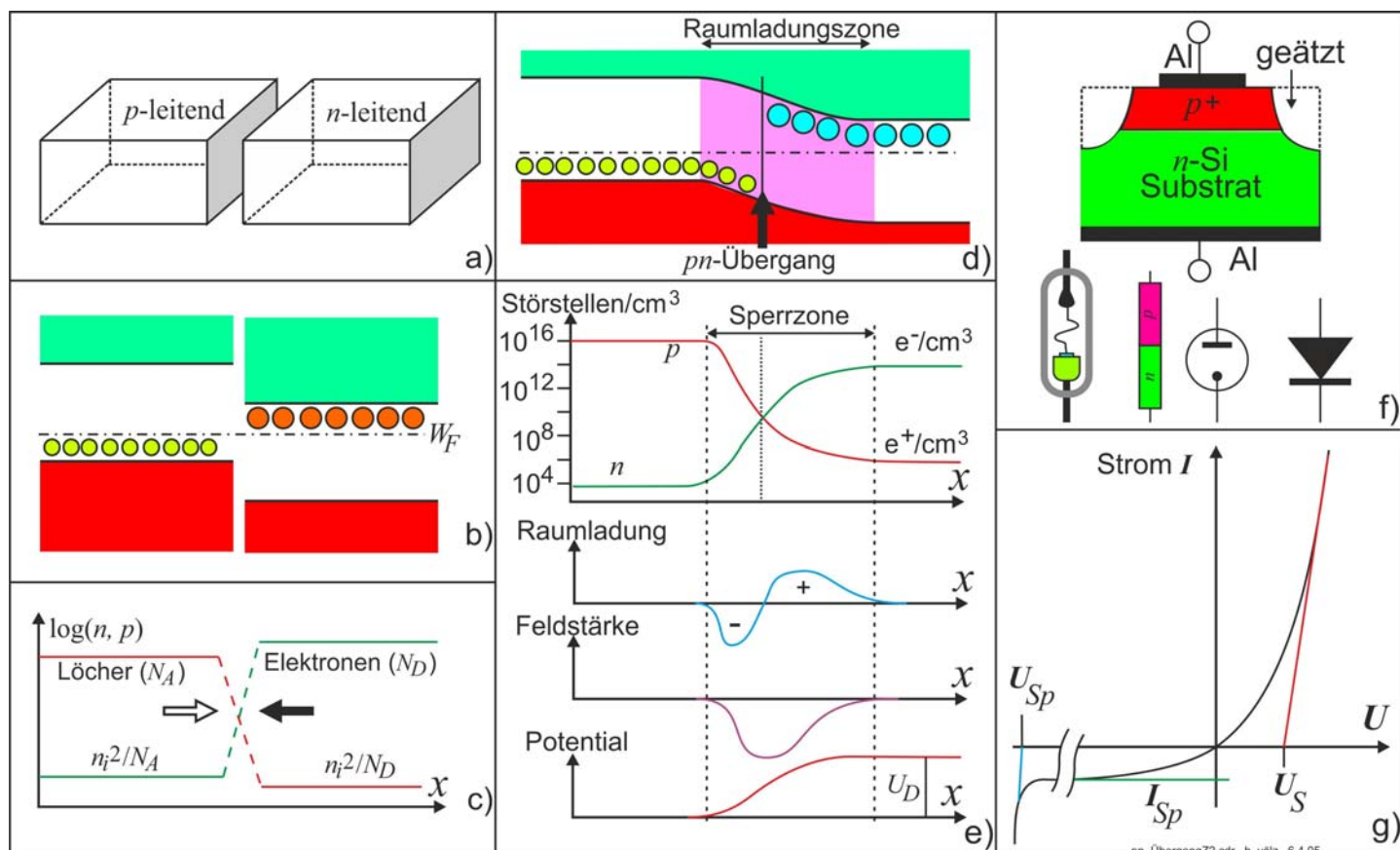


Bild 19. Eigenschaften und Kenndaten eines pn -Überganges.

Transistoren

Es sind mehrere Varianten von „verstärkenden“ Transistoren entstanden. Die wichtigsten Varianten sind in Bild 20 im Vergleich zur Elektronenröhre zusammengestellt. Grundsätzlich verschieden sind 1. der stromgesteuerte **Bipolar**-transistor und 2. der spannungsgesteuerte **Feldeffekt**-

transistor (FET). Beim FET wird die Leitfähigkeit eines Kanals verändert. Das erfolgt entweder durch eine Sperrschicht oder beim MOS-FET durch ein isoliertes Gate. Eine weitere Unterteilung für alle Transistoren geschieht durch die Verwendung bzw. Anordnung von p- und n-leitendem Material

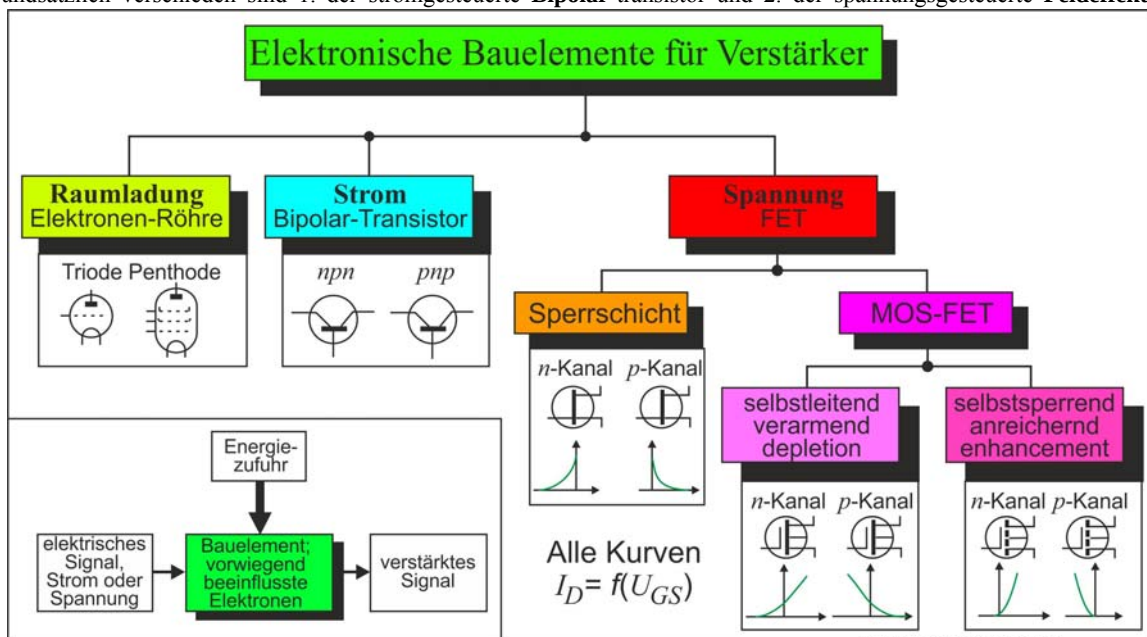
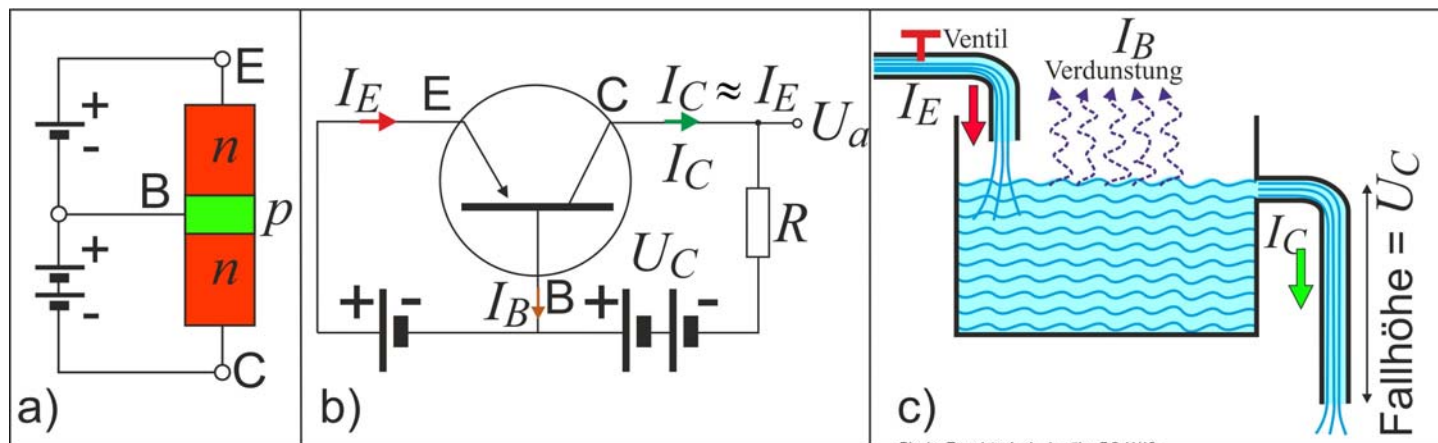


Bild 20. Die Varianten der Bauelemente für elektronischen Verstärker

Bipolar-Transistor

Der Bipolar-Transistor besteht im Wesentlichen aus zwei dicht aneinander gefügten pn -Übergängen (**Bild 21a**). Die gemeinsame mittlere p -Zone enthält so die beiden sich überlappenden Raumladungszonen. Das entsprechende Schaltbild zeigt (b).. Der mittlere Anschluss wird Basis B genannt. Bei einem Transistorbetrieb wird einer der pn -Übergang als **Emittor**⁸ E in Durchlassrichtung und der andere als **Kollektor**⁹ K in Sperrrichtung betrieben. Durch den Emittor-Strom werden Elektronen erzeugt und dann als Minoritätsträger in die gemeinsame Raumladezone injiziert¹⁰. Weil diese Zone technologisch kürzer als die typische Diffusionslänge der Elektronen gewählt wird, gelangen fast alle injizierten Elektronen in das Feld des anderen pn -Übergangs, den Kollektor, der in Sperrrichtung betrieben wird. In seinem Feld werden sie beschleunigt. Er „sammelt“ also fast alle injizierten Elektronen ein. Am Kollektor können sie als Strom abgenommen werden. Der Kollektorstrom I_C ist folglich fast so groß wie der eingespeiste Emittorstrom I_E . Da die Kollektorspannung U_{CB} jedoch erheblich größer als die Emitterspannung U_{EB} ist, erfolgt eine beachtliche Leistungsverstärkung. Die vergleichsweise sehr kleine Differenz $I_E - I_C$ bildet den Basisstrom I_B .

Die Vorgänge beim Bipolar-Transistor können durch (c) veranschaulicht werden. In das Wasserbecken fließt ein, mit dem roten Ventil einstellbarer Wasserfluss entspricht dem Emittor-Strom I_E . Die nahezu gleiche Menge fließt am Überlauf als I_C aus. Entsprechend der Fallhöhe U_C ist die



kinetische Energie am Ausfluss (Kollektor) entsprechend vergrößert. Die Differenz $I_E - I_C = I_B$ tritt in diesem Beispiel durch Verdunsten ein.

Bild 21. Zum Aufbau und zur Funktion des Bipolar-Transistors.

Neben den eben beschriebenen nnp -Transistoren gibt es auch **pnp-Transistoren**, wobei dann die Basis n -leitend, sowie Emittor und Kollektor p -leitend sind. Daher muss bei ihnen die Polarität der Spannungsquellen gewechselt werden. In die Raumladungszonen werden Löcher injiziert, deren Beweglichkeit deutlich geringer ist. Daher haben diese Transistoren meist schlechtere Eigenschaften und werden nur für Sonderzwecke eingesetzt.

Heute werden Transistoren überwiegend in **Planar-Technologie**¹¹ produziert. Dabei werden alle erforderlichen Halbleiterstrukturen und -anordnungen mittels lokal durchlässiger Masken erzeugt. Das geschieht durch Ätzen, epitaktisches¹² Aufwachsen, Oxidieren und Diffundieren. So entsteht eine Schichtenstruktur, wie sie **Bild 22a** schematisch zeigt.

Auf einem p -Substrat wird zunächst eine vergrabene n^+ -Schicht aufgebracht (blau). Darauf wächst epitaktisch die mit Antimon dotierte n -Schicht (rot). In sie wird durch eine Maske mit deutlich verkleinerter Fläche Bor als Akzeptor weniger hinein diffundiert. Infolge des Überwiegens der Bor-Konzentration entsteht dabei aus dem n -leitendem Material die p -leitende Basis (grün). Mit erneut verkleinerter Fläche und noch weniger Tiefe wird dann Phosphor diffundiert. Dabei entsteht durch Überwiegen des Donators Phosphor der n^+ -leitende Emittor (lila). Abschließend werden mit weiteren Masken die Al-Kontakte zu E, B und C hergestellt. Zum Schutz wird schließlich die Si-Oberfläche oxidiert. Einen Tiefenschnitt über die Mitte des Transistors für die so erhaltenen Dotierungen zeigt (b). Die vergrabene n^+ -Schicht dient vor allem zur elektrischen Trennung des Transistors gegenüber anderen Bauelementen auf dem gleichen Chip.

Vorwiegend wird der Transistor statt in der eingeführten Basisschaltung (c) die **Emitterschaltung** (d) benutzt. Wegen der Kirchhoff'schen Regeln gilt für die Ströme $I_E + I_B + I_C = 0$. Vom Basis- zum Kollektorstrom vermittelt dann eine 10- bis 1000-fache Stromverstärkung β gemäß

$$I_C \approx -\beta \cdot I_B.$$

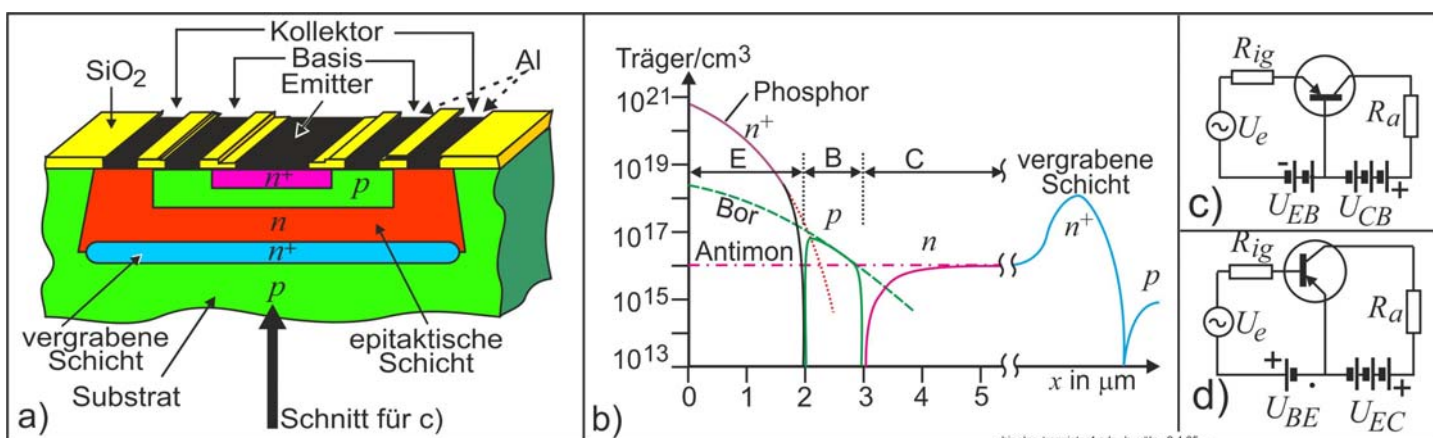


Bild 22. Aufbau und Schichtenstruktur des Planartransistors, sowie Basis- (c) und Emittor-Schaltung (d).

⁸Emittor lateinisch *emittere* herausgehen lassen, ausschicken.

⁹Kollektor lateinisch *collecta* Geldsammlung und *colligere* einsammeln.

¹⁰Injiziert lateinisch *inicare* hineinwerfen, einflößen,

¹¹Planar lateinisch *planarius* flach.

¹²Epitaxi griechisch *epi-* auf und *-taxia* (An)ordnen, taktos angeordnet,

Bei der Emitterschaltung ist die Stromverstärkung in einem beachtlich großen Arbeitsbereich nahezu konstant. Das hat u.a. Vorteile für die Kennliniendarstellung gemäß **Bild 23**. In den vier Quadranten werden auf den Ordinaten und Abszissen die verschiedenen Betriebsgrößen aufgetragen. Der I. Quadrant zeigt die Ausgangskennlinie $I_C(U_{CE})$ mit I_B als Parameter. Für sehr kleine Basisströme ist dies praktisch die Sperrkennlinie des Kollektor- pn -Übergangs (vgl. Gl. 1). Durch den Basisstrom wird sie nahezu parallel nach größeren Werten und ein klein wenig nach rechts verschoben. In einem gewissen Bereich ist dabei I_C nur sehr wenig von U_{CE} abhängig, das ist ähnlich wie bei der Penthode. Der II. Quadrant entspricht der obigen Stromverstärkung β . Die Kennlinie $I_C(I_B)$ verläuft weitgehend linear und ist nur wenig von U_{CE} abhängig. Der III. Quadrant zeigt etwa die Durchlasskennlinie des Emittor- pn -Übergangs, vgl. Bild 19g. Die Schwellenspannung beträgt in diesem Beispiel etwa 0,5 V. Die hohe Nichtlinearität dieser Kennlinien wird durch einen „Vorstrom“ $I_B \geq 40 \mu\text{A}$ und einen hinreichenden Vorwiderstand R_{ig} unterdrückt (Bild 21g). Der IV. Quadrant ist für die Anwendung wenig bedeutsam. Er zeigt den sich nur wenig ändernden Zusammenhang zwischen den beiden Spannungen U_{CE} und U_{BE} .

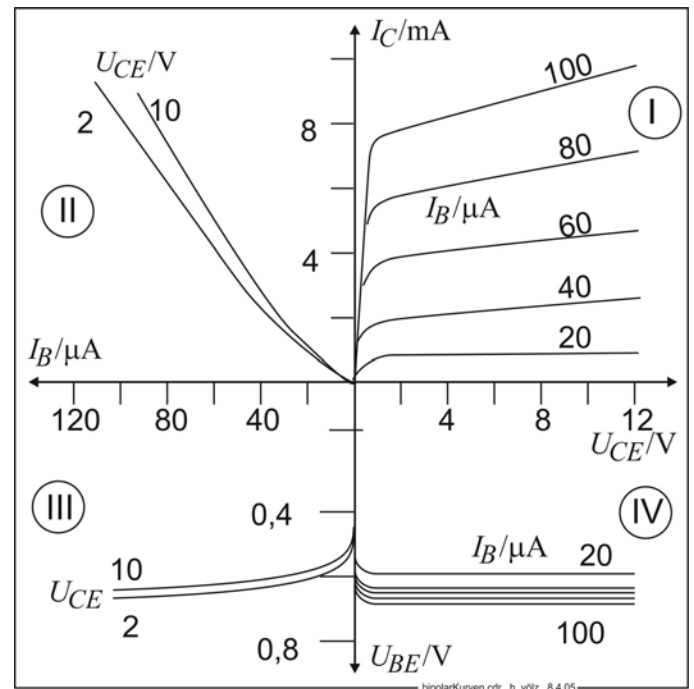


Bild 23. Kennlinienfeld eines Bipolartransistors in Emitterschaltung.

Sperrschicht-Feldeffekttransistor

Für die Steuerung des Ausgangsstromes mit Verstärkung benötigt der bipolare Transistor Strom und damit Leistung. Doch bereits 1924 ermöglichte die Elektronenröhre eine leistungslose Steuerung mit Spannung. Andererseits meldet Julius Edgar Lilienfeld bereits 1926 auf Grundlage der Drude-schen Elektronengastheorie das Prinzip eines FET (**F**eld-**E**ffekt-**T**ransistor) zum Patent an. Ein weiteres erfolgte 1934 von Oskar Heil. Aber erst fünfzehn Jahre nach dem bipolaren Transistor, nämlich um 1960 konnte ein MOS-FET technisch realisiert werden. Heute sind FET die wichtigsten Bauelemente der Mikroelektronik. Entsprechend Bild 19 gibt es heute mehrere Varianten.

Beim **Sperschicht-FET** (SGFET) wird die Raumladungszone des pn -Übergangs genutzt. Gebräuchlich ist auch der Name **JFET**¹³. Gemäß Bild 19e ist hier die Ladungsträgerdichte so gering, dass fast ein Isolator vorliegt, deshalb wird sie auch Sperschicht genannt. Ihre Breite hängt wesentlich von der Größe der angelegten Sperrspannung ab. Den historischen Aufbau eines SFET zeigt **Bild 24a**. Ein dünner Zylinder aus n -leitendem Halbleitermaterial ist an den Enden für den Stromfluss kontaktiert. Hierbei wird die linke Seite im Bild **Source**¹⁴ S genannt und dient der Einleitung des Stromes I_D . Dieser Anschluss ist auch das Bezugspotential für die weiteren Spannungen. Auf der rechten Seite wird der hindurch geführte Strom am Anschluss **Drain**¹⁵ D abgenommen. Zwischen beiden Anschlüssen entsteht die Spannung U_{DS} . Ein Teil des Zylinderumfangs ist p^+ -leitend dotiert (grün). Dieser Anschluss heißt **Gate**¹⁶ G. Hier wird eine Spannung in Sperrrichtung angelegt. Der dann notwendige Sperrstrom beträgt für Silizium weniger als 10^{-9} A. Dazu gehört dann ein Eingangswiderstand um $10^9 \Omega$. Änderungen der Gatespannung erfolgen daher nahezu leistungslos. Für die Funktion des SFET ist der entstehende Sperschichtbereich wesentlich. Seine Ausdehnung ändert sich mit der angelegten Spannung U_{GS} . Für die Source-Drain-Stromleitung steht daher nur ein Kanal mit dem Durchmesser d zur Verfügung. So wird der FET-Widerstand R_{SD} fast verlustfrei durch die Spannung U_{GS} gesteuert. Das gilt aber nur so lange, wie $U_{DS} \ll U_{GS}$ ist. Dann entstehen die Kennlinien (d). Die dortigen Bezeichnungen wurden von der Elektronenröhre übernommen. Mit U_{GS} ändert sich auch die wirksame Spannung zwischen p^+ -Gebiet und Kanal in Längsrichtung (c) und es nimmt die Breite der Sperschichtzone in Richtung Drain zu. Die Bilder (b) und (c) zeigen nur die obere Hälfte von (a). Mit steigendem U_{DS} erfolgt auf der Drain-Seite schließlich eine nahezu vollständige Abschnürung des Kanals. Dann liegt der Abschnür-, Penthoden-, Sättigungsbereich vor, was auch als Pinch-off-Effekt¹⁷ bezeichnet wird. Bei sehr großen U_{DS} entsteht im Material des Kanals eine Feldstärke, welche Ladungsträger generiert. Dadurch geschieht ein Durchbruch, – in (d) rot – mit einem steilen Stromanstieg. Da hierdurch der SFET zerstört werden kann, ist dieser Bereich im Betrieb zu vermeiden. Ideal sind alle geschilderten Eigenschaften beim planaren Aufbau (e) und zugleich mit Vertauschung von $p \leftrightarrow n$ -Leitung erfüllt.

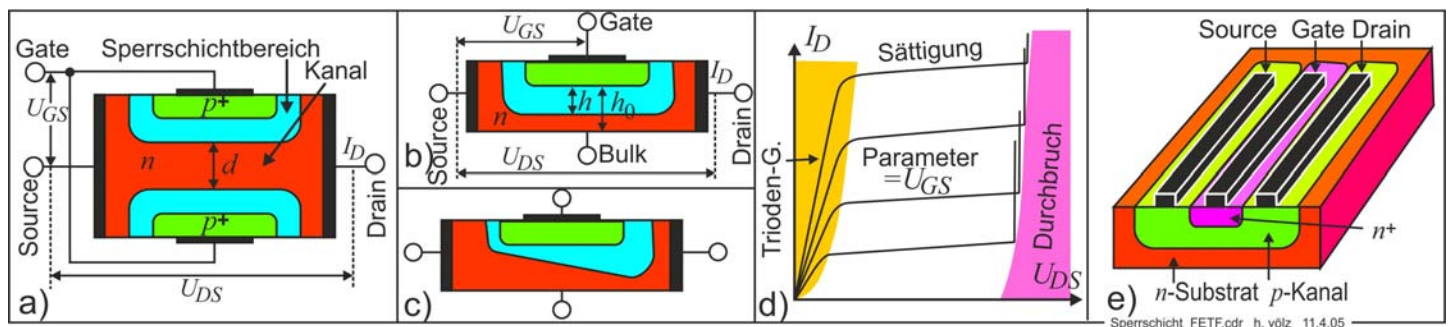


Bild 24. Vom historisch SGFET zum planaren Aufbau.

¹³ **J** nach *englisch* junction Anschluss, Knoten Übergang.

¹⁴ **S** Source *englisch* Quelle.

¹⁵ **D** Drain *englisch* Abfluss, Abzug, Senke.

¹⁶ **G** Gate; *englisch* Tor, Tür, Zugang, Eingang, Pforte.

¹⁷ Pinch-off *englisch* pinch quetschen, klemmen, einengen, drücken, beschränken, beengen, off weg, weg, weit, fort

MOS-FET

Eine Ausgangsstruktur für den MOS-FET (**metal oxide semiconductor**¹⁸) ist der **Metall-Isolator-Halbleiter-Übergang** (MIS). Er wird auch bei Umlaufspeichern und den CCD-Sensoren der digitalen Fotografie benutzt. In **Bild 25a** besteht aus einem n -leitenden Si-Kristall. Darüber befinden sich ein Isolator und dann der Metall-Kontakt. Beim MOS-FET (b) ist der der n -leitende Kristall ähnlich wie bei einem n p n -Bipolar-Transistor durch zwei gegeneinander geschaltete pn -Übergänge ersetzt. Jedoch gibt es zwei deutliche Unterschiede. Einmal sind die beiden pn -Übergänge so weit voneinander entfernt, dass keine gemeinsame Raumladungszone entstehen kann. Durch die große¹⁹ Entfernung ist sogar jede Wechselwirkung ausgeschlossen. Daher kann, wenn an den beiden n -Anschlüssen eine Spannung angelegt wird, zunächst kein Strom fließen. Denn je nach Polung der Spannung, ist immer einer der beiden pn -Übergänge gesperrt. Der zweite Unterschied besteht darin, dass das Gebiet zwischen den beiden pn -Übergängen gemäß einem MIS beeinflusst werden kann. Bereits 100 nm genügen für 100 V Durchschlagsfestigkeit, und der Isolationswiderstand beträgt ca. $10^{15} \Omega$. Je nach der am Gate angelegten Spannung tritt im gegenüber liegenden Abschnitt eine Ladungsträgerinversion als p -Kanal auf. So ergibt sich der Prinzipaufbau (b). Es gibt auch die komplementäre pnp -Variante. Bei ihr entsteht dann im Kanal ein p -Gebiet (p -Kanal-Typ).

Für den MOS-FET ergeben sich fast die gleichen Eigenschaften, Kennlinien usw. wie beim SGFET. Die Vorteile des MOS-FET bestehen vor allem darin, dass infolge des isolierten Gates kein temperaturabhängiger Sperrstrom auftritt und das Gate sowohl mit positiven Spannungen betrieben werden kann. Andere Bezeichnungen für MOS-FET sind MIS-FET und IGFET²⁰ (insulated gate; *lateinisch insularis* eine Insel, Inseln betreffend). Im Gegensatz zum Bipolar-Transistor heißt er auch Unipolar-Transistor, denn es kommt jeweils nur eine Trägerart vor: Elektronen beim n -Kanal-Typ und Löcher (Defektelektronen) beim p -Kanal-Typ. Eine Sonderform ist der TFT (Thin-Film-Transistor; Dünnschichttransistor). Er wird in dünnen, polykristallinen Schicht aufgedampft oder in Dickfilntechnik produziert.

Bei den Kennlinien eines MOS-FET treten im Gegensatz zum SGFT mehr Einflüsse auf:

- Bei der Herstellung der SiO_2 -Schicht durch Oxydation des Si wird im Isolator meist eine **positive Ladung** dauerhaft gebunden, was technologisch nur wenig zu beeinflussen ist.
- Metall und Halbleiter besitzen unterschiedliche **Austrittsarbeiten**, deren Differenz sowohl positiv als auch negativ sein kann. Durch Wahl der Materialien kann so die effektive Gate- und damit auch Schwellenspannung in beide Richtungen verlagert werden.
- An den Oberflächen des Halbleitermaterials können **Unregelmäßigkeiten des Kristallaufbaus** auftreten, die sich als Störterme auswirken,

Zusammen bewirken sie eine Verlagerung des Nullpunkts der $I_D(U_{GS})$ -Kennlinie. Für den n -Kanaltyp gilt Bild 11c. Es werden der **Anreicherungs-Typ** (A-Typ, c) und der **Verarmungs-Typ** (V-Typ, d) unterschieden. Je nachdem fließt bei $U_{GS}=0$ Strom oder nicht. Das Beispiel der vollständigen Kennlinie vom Anreicherungstyp eines MOD-FET zeigt e).

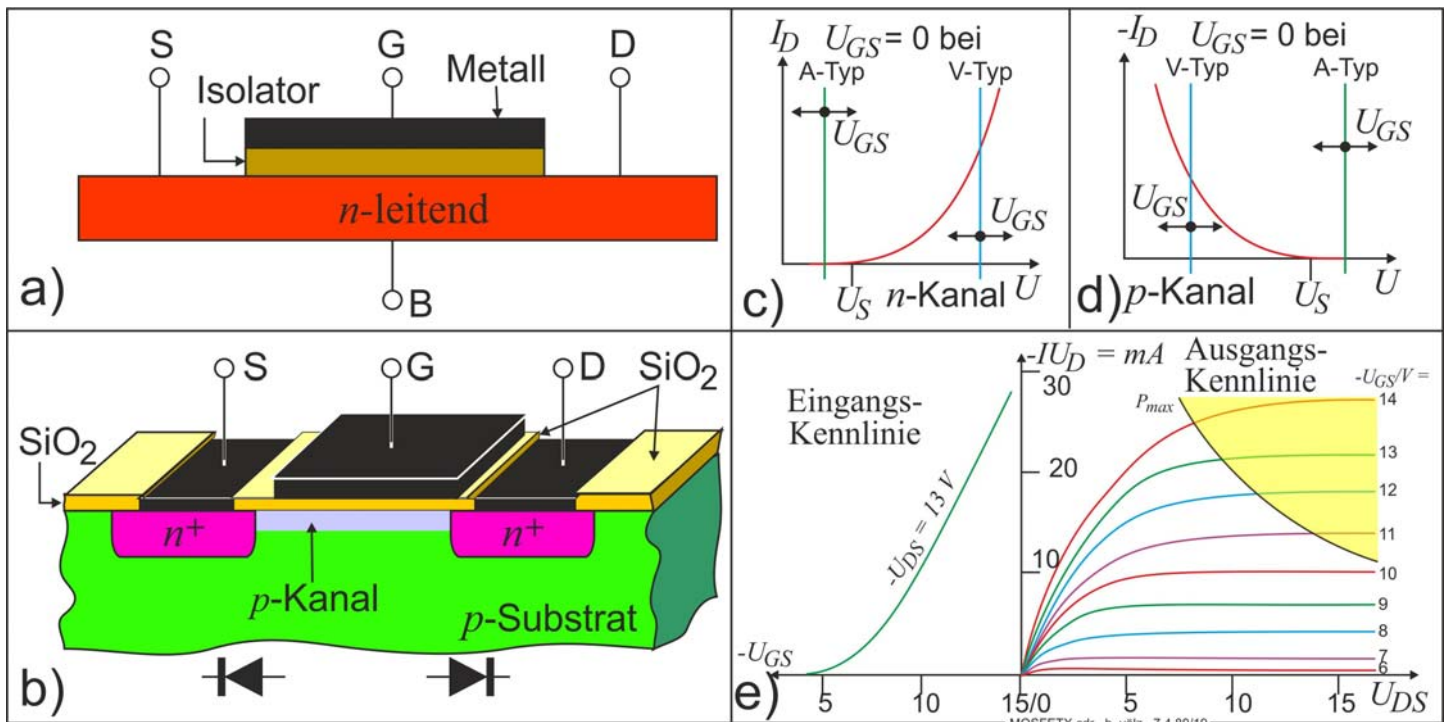
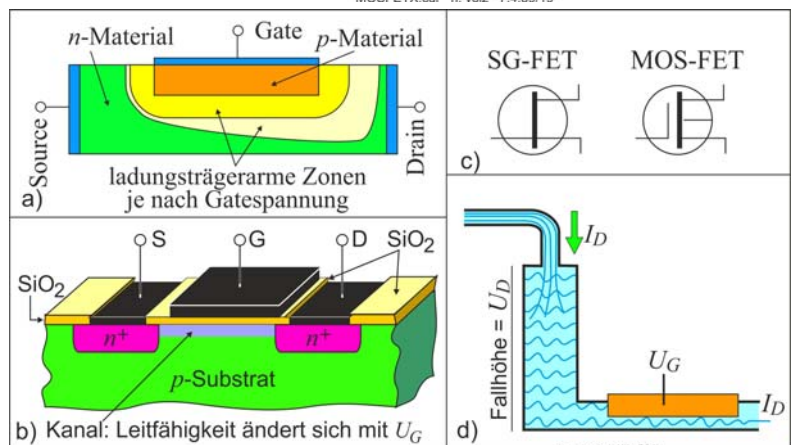


Bild 25. Aufbau und Wirkungsweise der MOS-FET.

Auch die Funktion der FET kann mittels der Analogie von **Bild 26** verständlich werden: a) gilt für den SG-FET und b) für den MOS-FET. Die Schaltbilder zeigen c) und d) demonstrieren das Analogie-Modell. Die Betriebsspannung des Kanals U_D bestimmt die Fallhöhe. Mit der Gate-Spannung U_G wird die Durchflussfähigkeit für den Strom I_D gesteuert.

Für höchste Frequenzen muss – analog zur Scheibentriode (Bild 16, S. 10) – der Kanal auf recht spezielle Weise stark verkürzt werden. So entsteht hier die komplizierte dreidimensionale Kurzkanaltechnik.

Bild 26. Analogie-Modell für die beiden FET.



¹⁸ *lateinisch conductor* Mieter, Pächter und *semi* halb, statt semiconductor wird zuweilen auch das S auf Silizium bezogen.

¹⁹ Das „groß“ ist hier relativ, denn für kurze Schaltzeiten muss die Kanal-Länge sogar sehr kurz, bis zu wenigen nm betragen. Daher gibt es Besonderheiten für sehr kurze Kanäle.

²⁰ *insulated lateinisch insularis* eine Insel, Inseln betreffend.

Transistor Geschichte

1925 JULIUS EDGAR LILIENTHAL (1881 - 1963) meldet elektronisches Bauelement (**FET**) an
 1934 OSKAR HEIL konstruiert und patentiert ersten **FET mit isoliertem Gate**
 1939 Konzept von Shockley für einen FET
 1948 JOHN R. PIERCE prägt **Begriff Transistor**
 1942 HERBERT MATARÉ experimentiert bei Telefunken mit **Duodiode** für Doppler-Funkmeß-Systeme (RADAR)
 1947 William Bradford SHOCKLEY (1910 - 1989) und John PEARSON (1908 - 1991) funktionierenden **Bipolar-Spitzen-Transistor** Bell Labs
 1948 Patentanmeldung für europäische Erfindung „**Transistron**“
 1955 Beginn der **Planartechnik**
 1956 SHOCKLEY, JOHN BARDEEN (1901 - 1991) und WALTER Houser BRATTAIN (1902 - 1987) Nobelpreis für Physik
 1960 erster MOSFET
 1962 P. WEIMER entwickelt **Dünnschichttransistoren** (thin film transistor, **TFT**)
 1966 CARVER MEAD entwickelt **Feldeffekttransistor** (MESFET) mittels **GaAs**

Betriebsarten der elektronischen Verstärker

Es sind zwei Verstärker-Anwendungen zu unterscheiden:

- Bei den **Signal-Verstärkern** sind die Ein- und Ausgangs-Signale *sehr klein* und damit gegenüber der notwendigen **Hilfsenergie** unwesentlich. Sie können daher unabhängig von der Hilfsenergie betrachtet werden. Sie existieren in vielen Varianten, u. a. mit Rückkopplungen, als Operationsverstärker mit verschiedenen Funktionen, als Frequenzumsetzer (Mischung) und besonders störarm.
- **Leistungs-Verstärker** sollen eine erhebliche **Energie des Ausgangssignals** bereitstellen. Im Vergleich sind daher andere Energie-**Verluste** vernachlässigbar. Weiter ist zu beachten, dass alle Energie-Verluste in **Wärme umgesetzt** werden, was eine zusätzliche **Kühlung** erfordert. Daher ist meist der **Wirkungsgrad** der Wandlung von der Hilfs- zur Ausgangsenergie wesentlich. Das erfordert spezielle **Schaltungen** und **Betriebsweisen** (A bis C). Im weitesten Sinne gehören hierzu auch die Regler- und Steuerungsschaltungen (s. ###).

Um den **Wirkungsgrad** der Umwandlung der Hilfsenergie $U_A \cdot I_A$ zur verstärkten Nutzenergie zu erhöhen, wird die einfache Verstärkerschaltung nach dem Prinzip von Bild 12 oder Bild 21d durch Gegentaktschaltungen ersetzt **Bild 27**, die gleichermaßen für Röhren und Transistoren gilt. Zunächst wird dadurch nur die Ausgangsleistung durch die zwei Endstufen verdoppelt. Bei (a) müssen die beiden hier gewählten MOS-FET mit gegenphasigen Spannungen U_g und $-U_g$ betrieben werden. Der in der Mitte angezapfte Transformator setzt dann die beiden ebenfalls gegenphasigen Ausgangsspannungen zusammen und transformiert sie unterteilt zum Lautsprecher. Infolge ihrer niedrigen Betriebsspannung können Transistoren aber problemlos auch in Reihe betrieben werden (b). Dennoch müssen sie gegenphasig angesteuert werden. Die dazu notwendigen beiden Steuerspannungen werden am Emitter bzw. Kollektor von T_0 abgeleitet.

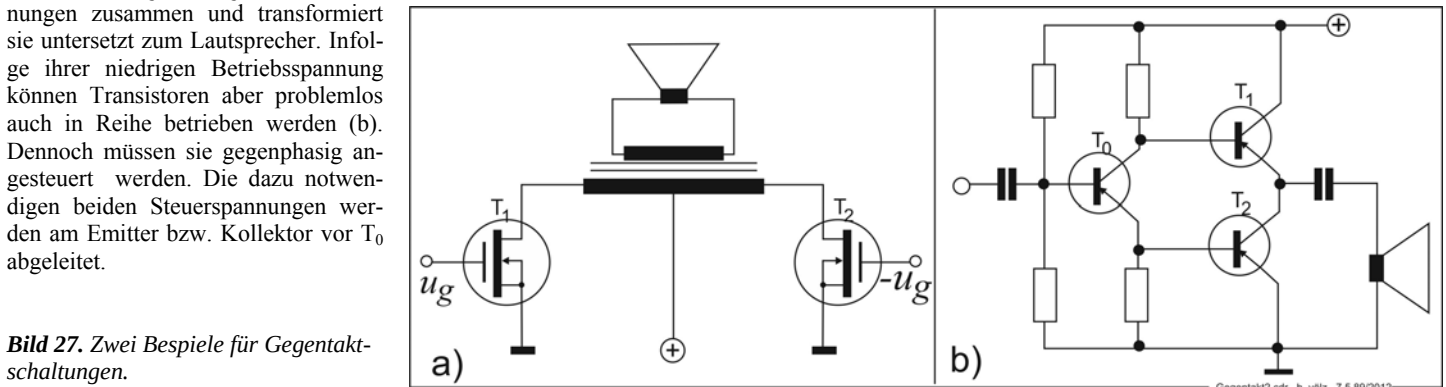


Bild 27. Zwei Beispiele für Gegentaktschaltungen.

Der entscheidende Gewinn der Gegentaktschaltungen wird jedoch durch die mögliche Verlagerung des Arbeitspunktes erreicht. Die verschiedenen Betriebsarten A bis C weist **Bild 28a** für die Eingangskennlinie aus. Der klassische **A-Betrieb** (d) benutzt nur den nahezu linearen Teil der Kennlinie und im Gegentakt werden einfach die beiden Ausgangsspannungen addiert. Als nützlichen Nebeneffekt tritt dabei eine deutliche Verringerung der immer etwas vorhandenen Nichtlinearitäten auf. Denn bei der Aussteuerung der treten bei beiden Transistoren gegenläufigen Krümmungen auf und so keine ungeradzahigen Oberellen entstehen. Das ermöglicht auch eine etwas weitere Aussteuerung und damit einen Gewinn an Ausgangsleistung. Jedoch der Gewinn wird erst wesentlich beim **AB-Betrieb** (b) erhöht. Ohne Eingangsspannungen ist der Ruhestrom in beiden Transistoren erheblich geringer und wegen der geringeren Verzerrungen ist auch eine deutlich weitere Ausnutzung der Kennlinien möglich. Natürlich ist dabei die Ausgangsspannung jedes einzelnen Transistors erheblich verzerrt, aber die Addition beider ergibt wieder ein sehr verzerrungsarmes Summensignal. Das ist beim **B-Betrieb** nicht mehr der Fall (e). Er hat daher nur noch spezielle Bedeutung und wurde meist gleich in den **C-Betrieb** (f) überführt. Dann treten am Ausgang nur noch kurzzeitige Impulse mit vielen geradzahigen Oberwellen auf. Das war früher bei den HF Sendern ohne Belang und ermöglichte so vergleichsweise hohe Sendeleistungen. Als Wirkungsgrad der Schaltungen ist das Verhältnis von Ausgangsleistung zur Leistung der Betriebsspannung eingeführt. In Abhängigkeit von der relativen Aussteuerung gilt dann (c).

Der **D-Betrieb** benutzt ein deutlich **anderes Verstärkerprinzip**, das schon sehr lange bekannt ist, aber erst mittels der **Digitaltechnik** technisch nutzbar wurde. Dabei steht D aber nicht für digital sondern ist eine schlichte Fortzählung von A über B und C. Dennoch ist meist der Begriff **D-Verstärker** üblich. Seine Wirkungsweise zeigt **Bild 29**. Es arbeitet mit möglichst exakten Rechteckschwingungen, wobei die Flanke zwischen der positiven und negativen Spannung systematisch verschoben wird (b). Mit der Lage der Flanken ändert sich der Gleichwert des Rechteckimpulses. Das führt zum Schaltbild (a). Das zu verstärkende Signal wird einem Komparator zugeführt, der auch mit einer periodischen Sägezahnsschwingung versorgt wird. Ihre Frequenz muss allerdings wesentlich höher als die höchste Signalfrequenz sein. So entstehen aus dem (in c grünen) Signal die schwarz gekennzeichneten Schaltpunkte und der Schaltverstärker erzeugt dann die (roten) Rechteckimpulse (c). Ihr Gleichwert wird durch den Tiefpass zurück gewonnen und entspricht dem verstärkten Eingangssignal. Da der Schaltverstärker nur Strom oder nicht Strom ausgibt, wird praktisch die gesamte Hilfs- bzw. Gleichleistung in Wechselleistung umgesetzt, was einem Wirkungsgrad von 100 % entspricht. Während beim AB-Betrieb mit integrierten Schaltkreisen kaum 10 Watt Ausgangsleistung zu erreichen sind, sind so erheblich größere Werte möglich. Außerdem hängt die Linearität des Ausgangssignals nur von der Linearität des Sägezahns ab.

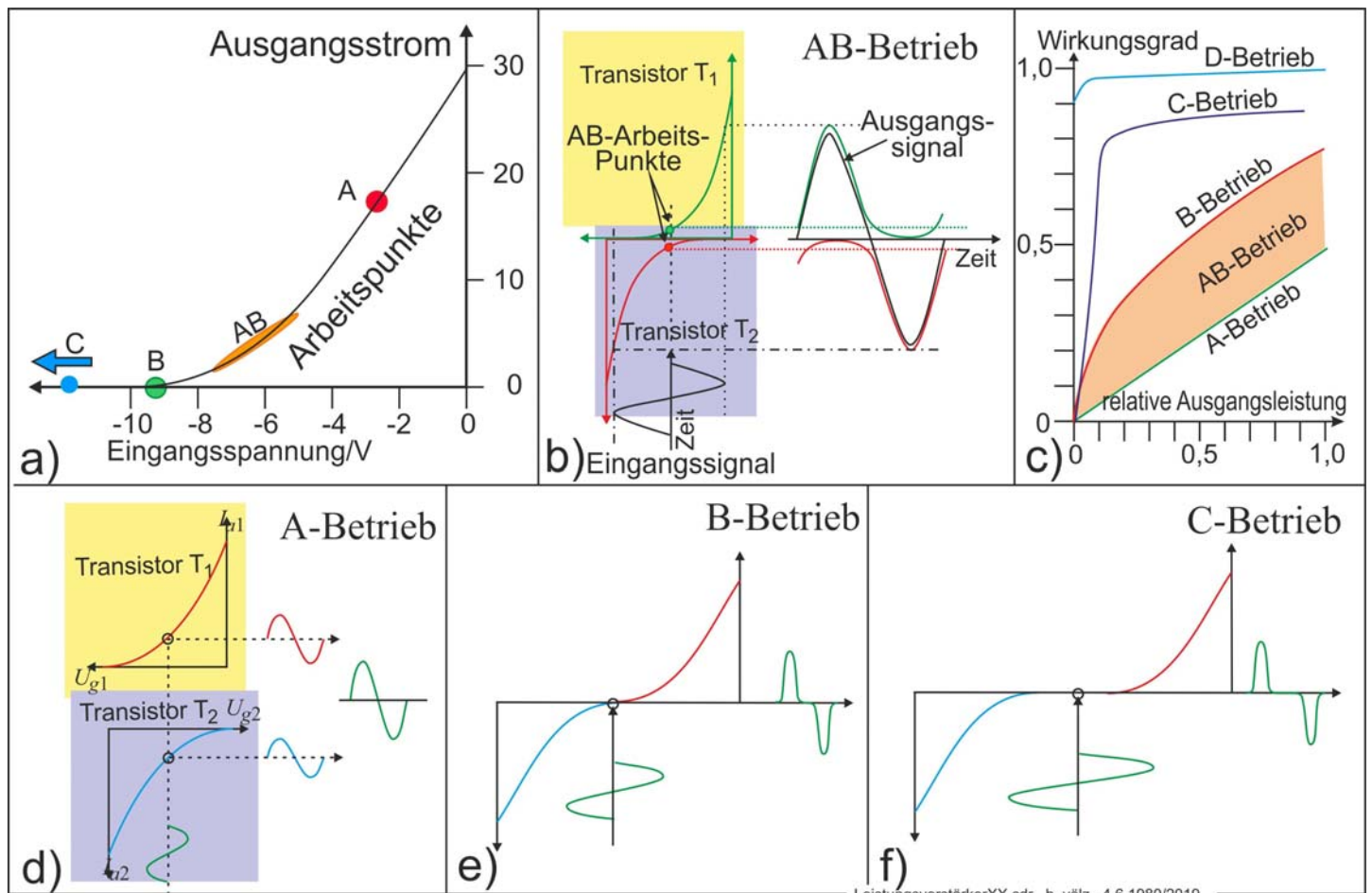


Bild 28. Die verschiedenen Betriebsarten von Gegentaktschaltungen.

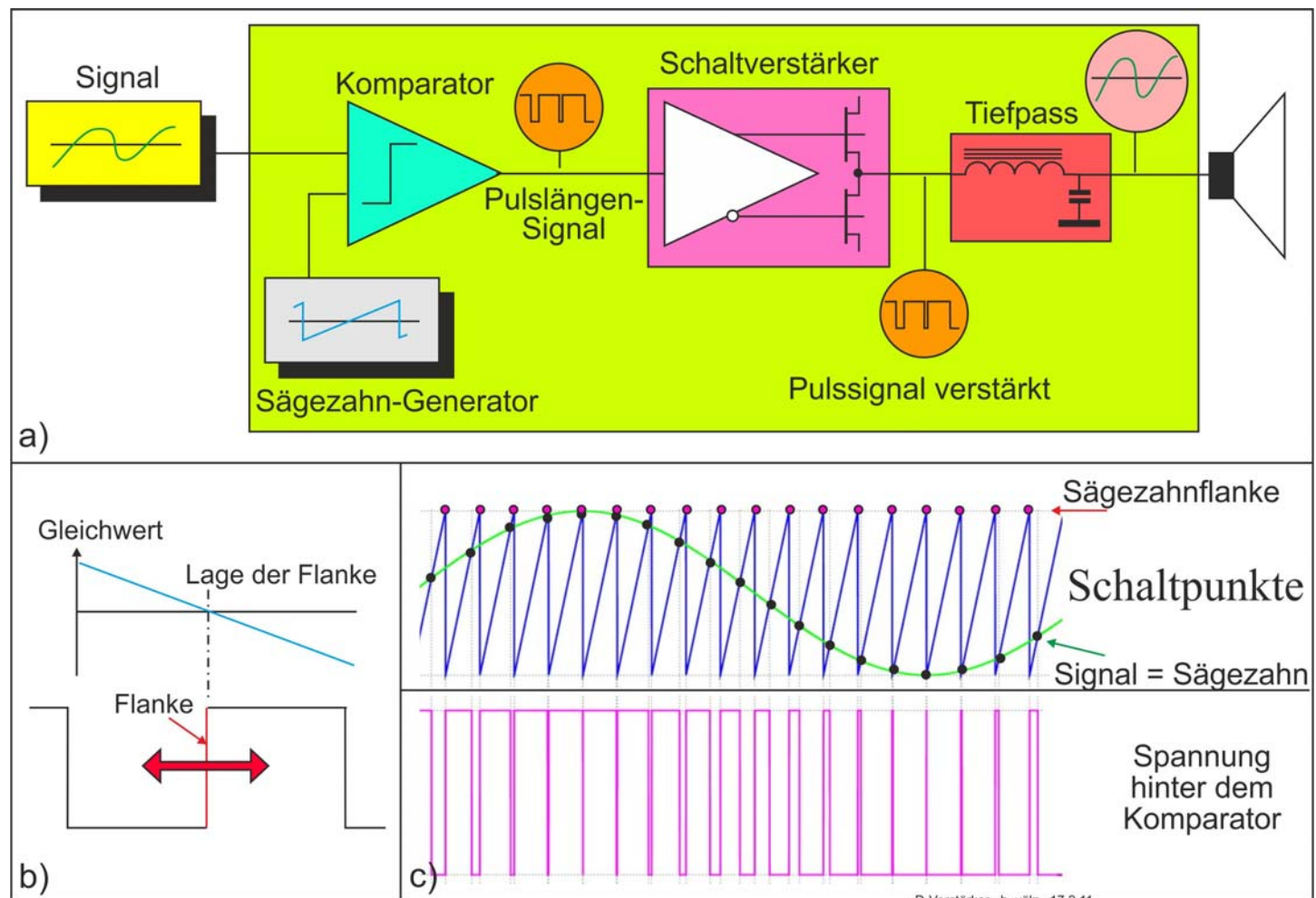


Bild 29. Prinzip des D-Verstärkers.

Es ist erstaunlich, dass die **Magnetbandaufzeichnung** eine beachtliche Ähnlichkeit mit den Betriebsarten der Gegentaktverstärker aufweist. Hierzu sei von der magnetische Hysteresekurve in **Bild 30a** ausgegangen. Bei den sehr frühen Aufzeichnungen wurde mit Gleichstromvormagnetisierung gearbeitet. Am einfachsten erfolgt sie mit der Neukurve. Dann ist nur der kleine lineare Bereich in (b) nutzbar. Er entspricht etwa dem A-Betrieb. Prinzipiell ist auch die Grenzhysterese mit einem weitaus umfangreicheren Nutzbereich möglich. Doch dann ist zuvor die Sättigung zu erzeugen und infolge der immer statistisch verteilten Magnetitpartikel tritt ein sehr starkes Rauschen auf. Die hervorragende Qualität der HF-Vormagnetisierung stellt eine gewisse Analogie zur Gegentaktschaltung dar. Mit den beiden Spitzenwerten der Hochfrequenz werden die beiden verschobenen Neukurven erzeugt (d). Sie entsprechen den beiden gegeneinander geschalteten Transistoren mit dem Unterschied dass sie aber nacheinander eingestellt werden. Im zeitlichen Mittel erzeugen sie gemittelt die sehr gut lineare Remanenzkurve. Aus der HF-überlagerten Signalspannung (in c rot) entsteht dann die erheblich verformte Kurve (e). In folge der nicht wiedergebbaren HF bleibt dabei aber nur das unverzerrte Signal (in e schwarz) übrig. Für dieses Ergebnis muss jedoch die HF-Amplitude die richtige Größe in Bezug auf die Hysterese des jeweiligen Magnetbandes besitzen. Die weist (f) aus. In der Praxis wird meist der Wert des zweiten Klirrfaktorminimums k_{min2} benutzt.

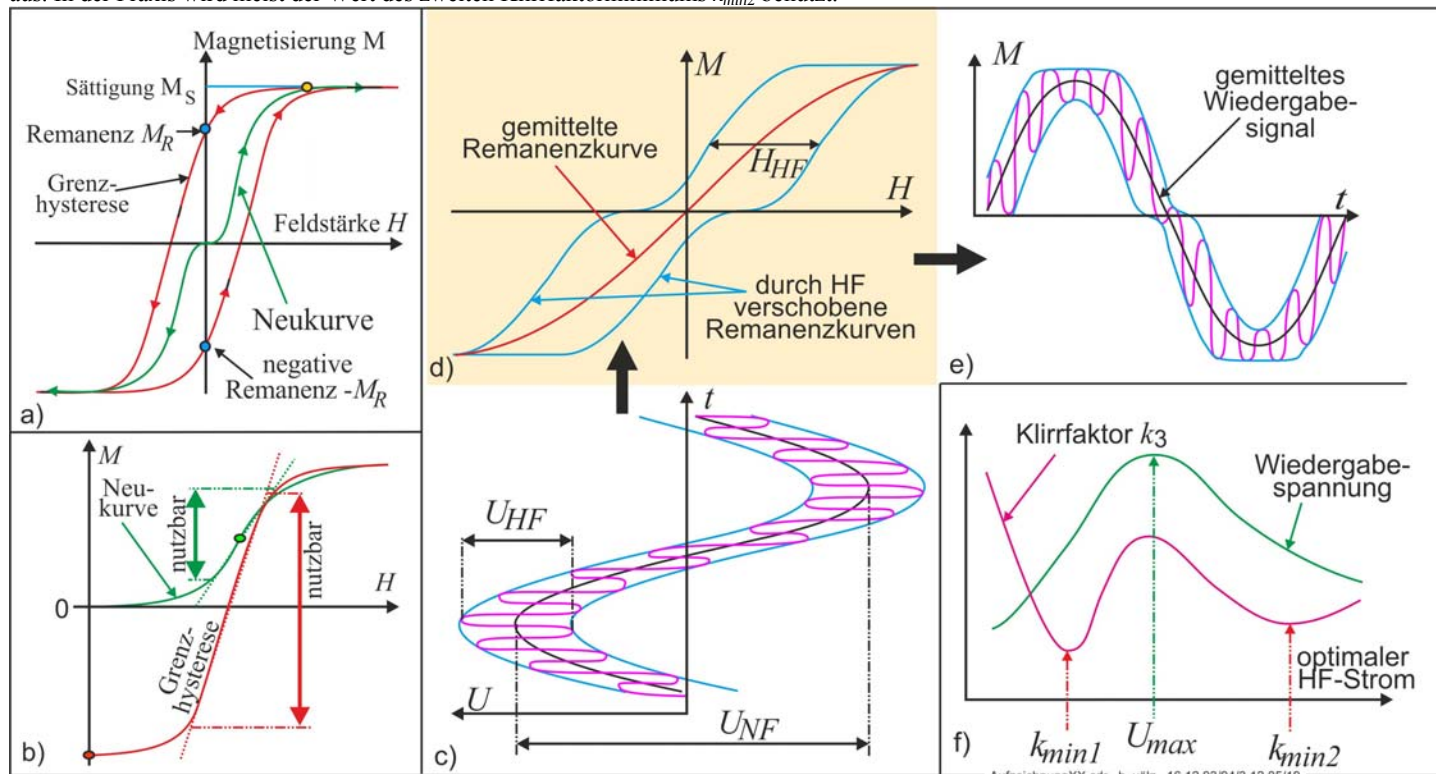


Bild 30. Zur Analogie der HF-Vormagnetisierung beim Magnetband mit den Gegentaktschaltungen.

Rückkopplungen

Der Begriff wurde wohl erstmalig 1904 von Lord Kelvin of Largs (William Thomson) bei Röhrenverstärkern benutzt (Englisch feed back). Ab etwa 1912 wurde sein Inhalt dann mehrfach betont u. a. ab 1912 von de Forest, Armstrong (USA) und Alexander Meißner. Jedoch viel früher wurde sie schon (ohne den Begriff zu verwenden) benutzt. Bereits in der Antike war die so mögliche Stabilisierung bekannt, z. B. ab etwa 300 v. Chr. für die Regelung der Öllampe oder für eine konstante Durchflussmenge durch das Schwimmventil. Genau so hat das Prinzip dann 1769 James Watt für seinen Fliehkraftregler benutzt (s. S. 22). Allgemein führte die Rückkopplung aber erst etwa 1940 Norbert Wiener bei seiner Kybernetik ein. Heute werden darunter viele Inhalte eingeordnet und so erlangt sie eine fundamentale Bedeutung und allgemeine Anwendung: Von der Regelung zur Stabilisierung ausgewählter Parameter, wie elektrische Spannung oder Temperatur, über die Kettenreaktion (Atombombe), Vermehrung und (berentzes) Wachstum (Evolution, Biologie), beim Geld (Zinseszins, Börse), als Rekursion und Iteration (Informatik, Fraktale, L-Systeme) und beim Verhalten (Psychologie) bis hin zur Soziologie, Kultur (Kommunikation, Meme, Gebräuche, Sitten).

Die genauere Betrachtung bei den elektronischen Verstärkern lässt besonders einfach für die Elektronenröhre beschreiben: Das erfolgt mit ihren drei Kenngrößen Steilheit S , Durchgriff D und Innenwiderstand R_i :

$$S = \frac{\partial I_a}{\partial U_g}; D = \frac{\partial U_g}{\partial U_a}; R_i = \frac{\partial U_a}{\partial I_a} \text{ sowie die Barkhausengleichung } S \cdot D \cdot R_i = 1.$$

Mit dem Arbeitswiderstand R_a folgt die Steuergleichung

$$i_a = S(u_i + D \cdot u_a) \text{ mit } u_a = -i_a \cdot R_a.$$

Für die Verstärkung gilt dann

$$V = \frac{u_a}{u_i} = \frac{1}{D} \cdot \frac{R_a}{R_a + R_i} = S \cdot \frac{R_a \cdot R_i}{R_a + R_i}$$

Bei $R_a \gg R_i$ folgt die (maximale) **Grenzverstärkung** μ

$$V_\infty = \mu = 1/D$$

1925 wurde so erstmals eine Verstärkung von etwa 100 000-fach mit mehreren Röhren erreicht.

Das Blockschaltbild für eine allgemeine Rückkopplung zeigt **Bild 31**. Dabei das Eingangssignal u_e zum Ausgangssignal u_a verstärkt und irgendwie verändert als u_r zum Eingang zurückgeführt:

$$u_r = -k \cdot u_a + l \cdot i_a \text{ und } \begin{cases} \text{sign } k \\ \text{sign } l \end{cases} = \begin{cases} + \\ - \end{cases} \text{ bei } \begin{cases} \text{Mitkopplung} \\ \text{Gegenkopplung} \end{cases}.$$

Es werden als unterschieden: Spannungs- oder Stromrückführung durch k bzw. l (Dimension ist R) und Mit- oder Gegenkopplung durch $+$ oder $-$ erfasst.

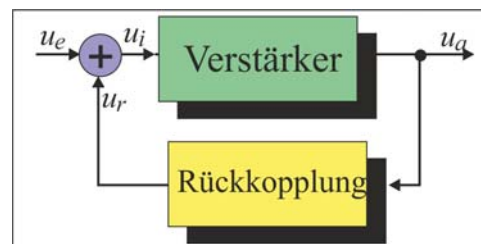


Bild 31. Blockschaltbild der Rückkopplung.

Alle vier Fälle erfolgen ohne Phasenverschiebung für die Rückführung und für die Kennwerte gilt dann:

$$19 \text{ und } S^+ = \frac{S}{1-l \cdot S}; \quad D^+ = D-k; \quad R_i^+ = \frac{1}{S^+ \cdot D^+} = R_i \cdot \frac{1-l \cdot S}{1-k}$$

Folglich ändert eine Spannungsrückkopplung nur den Durchgriff D und die Stromrückkopplung nur die Steilheit S . Beide wirken sich aber auf den Innenwiderstand R_i aus. Die Formeln gelten für jeden beliebigen (linear angenommenen) Verstärker unabhängig von der Stufenzahl und der Phasendrehung zwischen Ein- und Ausgang, wobei dann allerdings die Steilheit S und der Durchgriff D eventuell auch negatives Vorzeichen besitzen können.

In Erweiterung zum Bild 31 zeigt eine mehr detaillierte Schaltungstechnik von **Bild 32**. Ohne auf die Formeln Bezug zu nehmen ist sie dann leicht auf Transistoren zu übertragen.

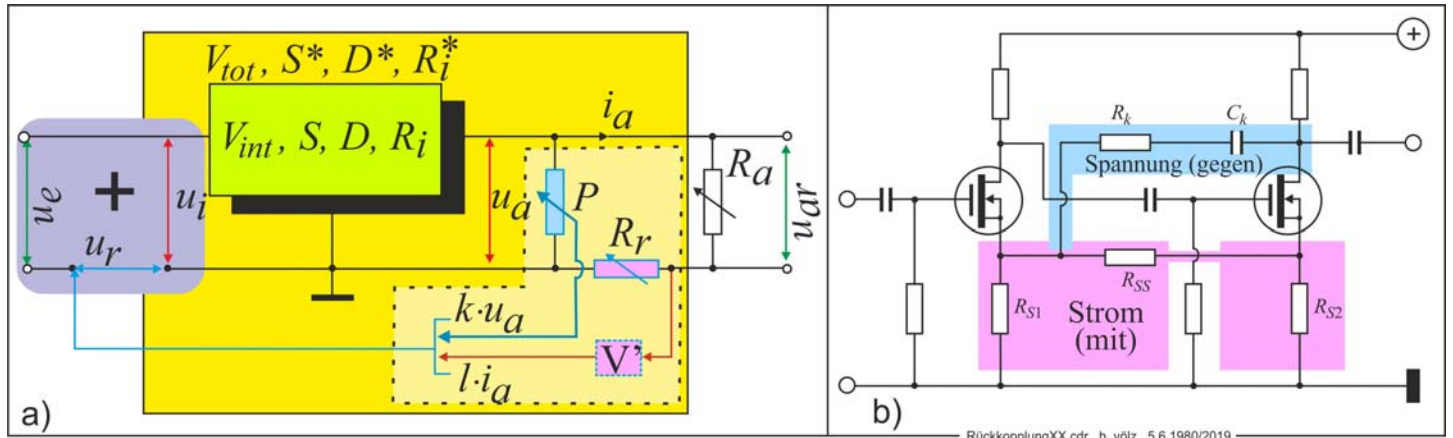


Bild 32. a) Prinzip der gemischten Rückkopplung in Erweiterung zu Bild 31, b) eine Ausführung mit MOS-FET.

Eine **Gegenkopplung** wirkt der Eingangsspannung entgegen und reduziert so die Verstärkung, wirkt meist **stabilisierend** auf alle Systemeigenschaften. Oft ist eine **Linearisierung** der Kennlinien besonders wichtig. Dann treten geringere Verzerrungen auf. Eine **Mitkopplung** wird zur Eingangsspannung addiert. Sie **erhöht die Verstärkung**. Durch eine **gemischte Gegen- und Mitkopplung** können ideale Spannungs- oder Stromquellen erzeugt werden. Doch wichtiger ist die Erzeugung extrem kleiner oder großer und sogar negativer **Innenwiderstände**. Ein konkretes Beispiel hierfür zeigt (b). **Negative Widerstände** sind nur mittelbar vorhanden, vgl. [Völ59]. Allerdings kommen sie mittelbar bei den Kennlinien von u. a. Tunnel-dioden, Gunn-Elemente und Gasentladungen vor. Dort gilt dann $U \uparrow \Rightarrow I \downarrow$ und umgekehrt und somit $R = -\Delta U / \Delta I$. Die mit gemischter Rückkopplung erzeugten negativen Widerstände wirken sich vor allem dadurch aus, dass sie einen Teil von R_a kompensieren. Das ist z. B. zur Bedämpfung der Eigenschwingungen von Lautsprechern sehr nützlich. Für Leitungen (**Bild 33**) können sie durch einfache Parallelschaltung des Ausgangs ohne Eingangsspannung deren Dämpfung senken (c) und so mittelbar wie zwei zwischengeschaltete Verstärker wirken (b) und dabei dennoch für beide Übertragungsrichtungen verstärkend wirken. So kann im Prinzip die sonst notwendige Gabelschaltung entfallen (a, b).

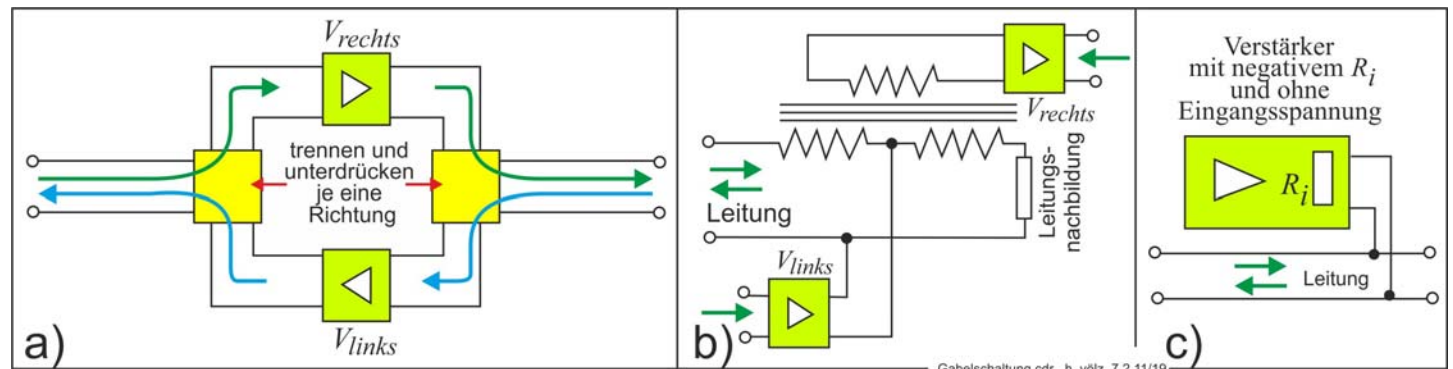


Bild 33. a) Prinzip der Gabelschaltung, (b) eine Ausführung mit Symmetrietrafo und (c) mit einem Verstärker mit negativem R_i .

Schließlich gibt es auch **Blindrückkopplungen**, die in den Rückführungen Phasendrehungen benutzen. So lassen sich u. a. steuerbare und/oder große Blindwiderstände als virtuelle Kapazitäten oder Induktivitäten erzeugen.

Bei **Impulsen** wirken Rückkopplungen auf spezielle Weise. Sein Wert N_0 wird zunächst um einen Faktor α zum Ausgang erhöht. Dann kehrt erneut über die Rückkopplung zum Eingang zurück. Das kann sich ständig fortsetzen und mit fortlaufender Zählung k gilt dann immer

$$N_{k+1} = \alpha \cdot N_k$$

So ergibt sich für die angegebenen speziellen Werte **Bild 34**. Es entsteht also ein diskret exponentieller Anstieg. (a). Dazu gehört mittelbar eine kontinuierliche Differentialgleichung

$$\frac{dN}{dt} = \alpha \cdot N(t)$$

Mit den Startwerten t_0 und N_0 führt die Integration exakt zu einem exponentiellen Wachstum:

$$N = N_0 \cdot e^{\alpha \cdot (t - t_0)}$$

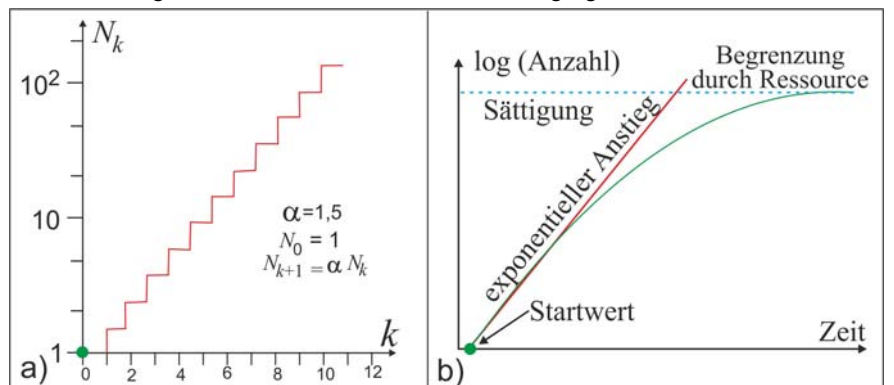


Bild 34. Exponentieller Anstieg durch sich wiederholender Rückkopplung und b) unter Einbeziehung von Begrenzung durch endliche Ressourcen

Durch Ressourcen-Grenzen mit einem Maximum M muss die Differentialgleichung abgewandelt werden:

$$\frac{dN}{dt} = \alpha \cdot N(t) \cdot \frac{M - N(t)}{M}$$

Dabei wird der Anstieg fortlaufend immer stärker reduziert und es tritt eine Sättigung auf (b)

Zur sehr großen Erhöhung der Empfindlichkeit von Audioschaltungen wurde das Iterationsprinzip 1920 von Edwin Armstrong bei seiner Pendelrückkopplungsschaltung – auch als Pendel-Audion und Super-Regenerativ-Empfänger genannt – erstmalig erfolgreich genutzt (**Bild 35**). Mit nur einer Röhre (bei der 1950 erfolgende Einführung des UKW auch mit Transistoren) waren so Verstärkungen von vielen Tausend möglich. In der Schaltung erfolgt die Rückkopplung vom Drain des Transistors auf die Hochfrequenzspulen, die mit dem Drehkondensator auf den zu empfangenden Sender abgestimmt werden. Zum Gate wird außerdem eine periodische Rechteckschwingung von >20 kHz eingefügt (c). Sie schaltet die Verstärkung des Transistors immer nur während der negativen Hallwelle ein. Vor dem Einschalten liegt nur eine relativ kleine Senderspannung vor (b). Sie wird dann exponentiell fortlaufend verstärkt, und erreicht am Ende der negativen Rechteckspannung ihr Maximum. Dann wird die Rechteckspannung positiv und so die Verstärkung sowie der exponentielle Anstieg beendet. So entstehen die exponentiellen Einzelverläufe (d), die proportional zur jeweiligen Sendeamplitude beim Beginn der negativen Rechteckflanke sind (e). Die Wirkung des Teils der Audioschaltung und durch den Kondensator am Ausgang wird die HF dort entfernt und es verbleibt nur die aufmodulierte NF (f).

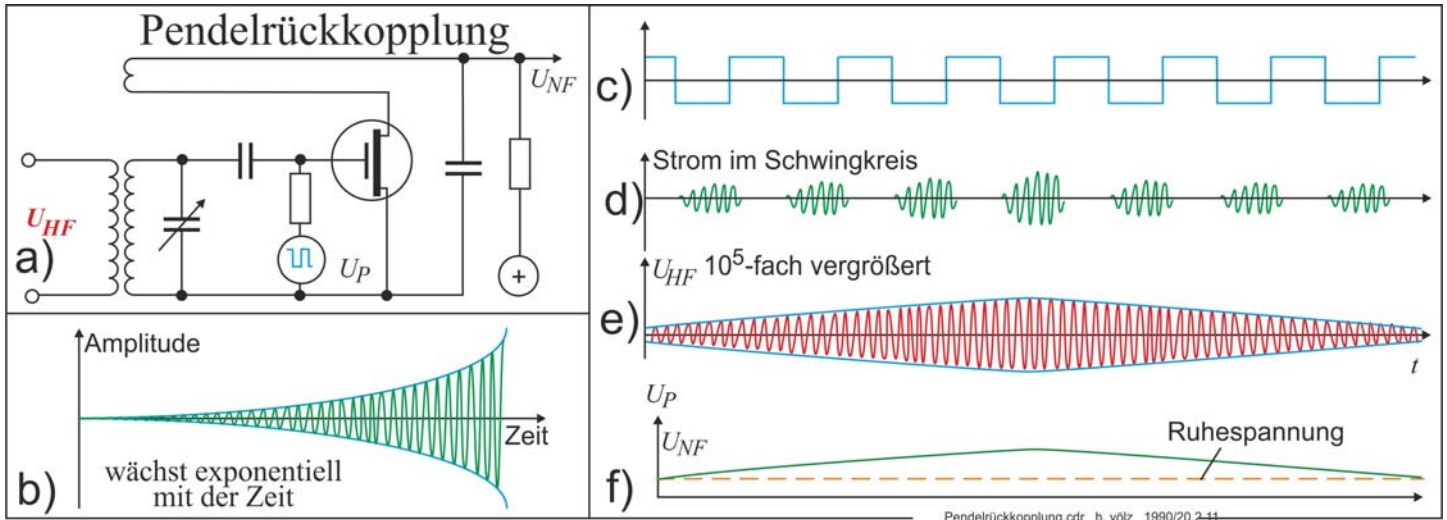


Bild 35. Prinzip und Wirkungsweise der Pendelrückkopplung.

Operationsverstärker

Sie sind vor allem für die Analogrechner entstanden, mit denen Differentialgleichungen simuliert werden können. Schematisch werden sie gemäß **Bild 36a** dargestellt. Sie besitzen dann zwei Eingänge, die mit + und - gekennzeichnet sind. mit der Verstärkung V bewirken die Eingangsspannungen U_{e1} und U_{e2} die Ausgangsspannung

$$U_a = V \cdot (u_{e2} - U_{e1}) \text{ bzw. } U_a = V \cdot U_e$$

Eine typische Eingangsstufe für Operationsverstärker (OV) zeigt (b). Der Strom I der Konstantstromquelle wird je nach den Steuerspannungen U_{e1} und U_{e2} auf die beiden Transistoren aufgeteilt und führt zu den unterschiedlichen Spannungen $-U_1$ und U_2 . Im OV wird dann nur U_2 zur Ausgangsspannung U_a weiter verstärkt. Ein OV ist also intern meist recht komplex aufgebaut. Viele weitere Details enthält u. a. [Völ89], S. 573ff. Im Bild 34 sind als Baustufen für Analogrechner vier Beispiele gezeigt. Den typischen Umkehrverstärker, dessen Verstärkungsgrad nur durch R_1 und R_2 bestimmt ist, zeigt (c). Er kann durch mehrere Eingänge zum Umkehr-Addierer erweitert werden (d). Dem Prinzipaufbau für einen Integrierer zeigt (e) mit der dazugehörigen mathematischen Beziehung. Alle Signalkomponenten müssen dabei oberhalb der angegebenen unteren Grenzfrequenz liegen. Den am meisten benutzten Differenzierer zeigt analog (f). Sie arbeiten besonders exakt für sehr kleine U_a .

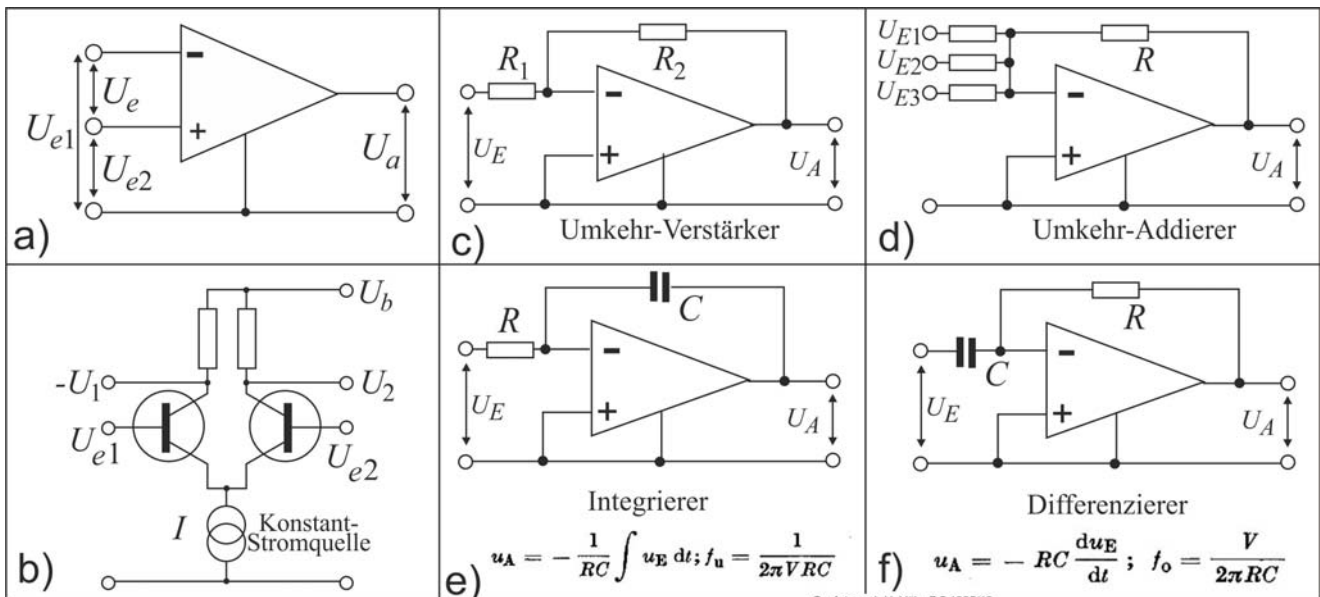


Bild 36. Aufbaudetails und wichtige Beispiele für Anwendungen der Operationsverstärker.

Neben den gezeigten Beispielen gibt es noch viele weitere Anwendungen. In [Völ89] sind über zehn 12 Varianten zusammengestellt, darunter Strom-Spannungs-Konverter, Spannungs-Strom-Konverter, Strom-Messer, Lineares Ohmmeter, Elektrometer-Verstärker, Elektrometer-Impedanzwandler sowie logarithmisch wirkende Schaltungen für Strom oder Spannung. Besonders vorteilhaft ist bei ihnen immer der konstant einstellbare Verstärkungsfaktor.

In Ergänzung zum negativen Ausgangswiderstand seien nur noch die mit dem OV mittelbar erzeugbaren Bildwiderstände ergänzt. Sie treten unterschiedlich am Eingang und Ausgang auf (**Bild 37**). Vor allen werden die Ausgangswiderstände genutzt. Für sie gilt (a) und (c) zeigt, dass mit R und C sowohl virtuelle Induktivitäten L als auch Kapazitäten C erzeugt werden können. Ihre Größe und Vorzeichen können über die Steilheit S entsprechend der Stärke und dem Vorzeichen der Rückkopplung verändert werden. Sie besitzen dabei auch immer einen gewissen Realanteil, der hier nicht angegeben ist. Für den weniger wichtigen Eingangswiderstand gilt (b) und

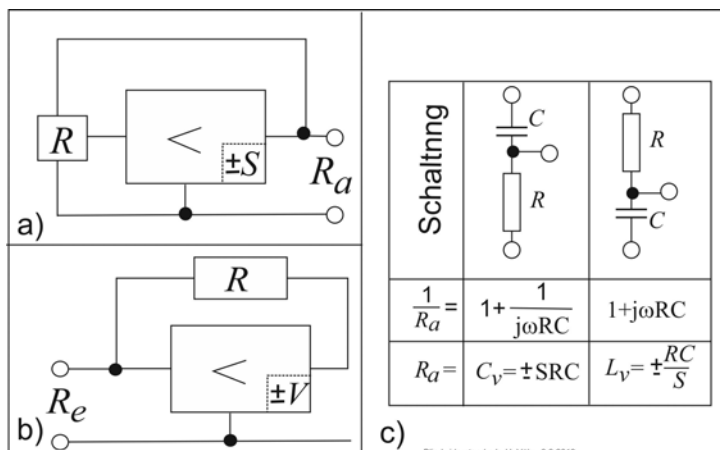


Bild 37 Zum Entstehen komplexer Ein- und Ausgangswiderstände.

$$R_e = R \frac{1}{1 \pm V} \text{ und speziell } R_{e,L} = \frac{j\omega L}{1 \mp V} \text{ bzw. } R_{e,C} = \frac{1}{j\omega C(1 \pm V)}$$

Weitere Details enthält [Völ89].

RC-Verstärker und Generatoren

Spulen (Induktivitäten) sind besonders kostspielige Bauelemente. Insbesondere in integrierte Schaltungen lassen sie sich nur für höchste Frequenzen herstellen. Das war ab etwa 1950 ein wichtiger Grund, frequenzselektive Schaltungen nur mit Widerständen und Kondensatoren herzustellen. So entstanden RC-Verstärker und -Generatoren²¹. Wesentlich sind dabei Zusammenschaltungen aus RC-Hoch- und -Tiefpässen. Von den vielen Möglichkeiten zeigt in **Bild 38a** die gelb unterlegte Schaltung ein besonders häufig benutztes Beispiel. Es ist hier in den Rückkopplungszweig eingeschaltet. Für das Übertragungsmaß $\ddot{u}_f$ der Schaltung aus R_1, R_2, C_1, C_2 gilt dann der Kreis in (b). Der Frequenzgang wird durch die drei am Verstärkereingang vorhandenen Spannungen U_e, U_i und U_r bestimmt. Für die selektive Verstärkung gilt weiter $V = U_a/U_e$. Anschaulich ist die Wirkungsweise nur der Schaltung schwer aus (b) abzuleiten, denn dazu muss der Betrag des Verhältnisse $|U_a/U_e|$ bestimmt werden. Für unterschiedliche Verstärkungen V ergeben sich dann die Kurven (c). Beim Erreichen einer $V_{osz} = 1/\ddot{u}_{f\text{maximum}}$ geht der RC-Verstärker in einen Oszillator²² für diese Frequenz über. Dies zu erklären bereitete zunächst Vielen große Schwierigkeiten. Denn es gab keinen Resonanzkreis, bei dem die Energie zwischen C und L wechseln konnte. So wurden zuweilen sogar virtuelle Induktivitäten in Betracht gezogen. Aber schließlich wurde klar, dass dann faktisch nur das thermische Rauschen dieser Frequenz verstärkt wurde. Deshalb schwankte auch die erzeugte Frequenz immer ein wenig. Weitere Details zu den RC-Verstärkern enthält wohl erstmalig [Völ56].

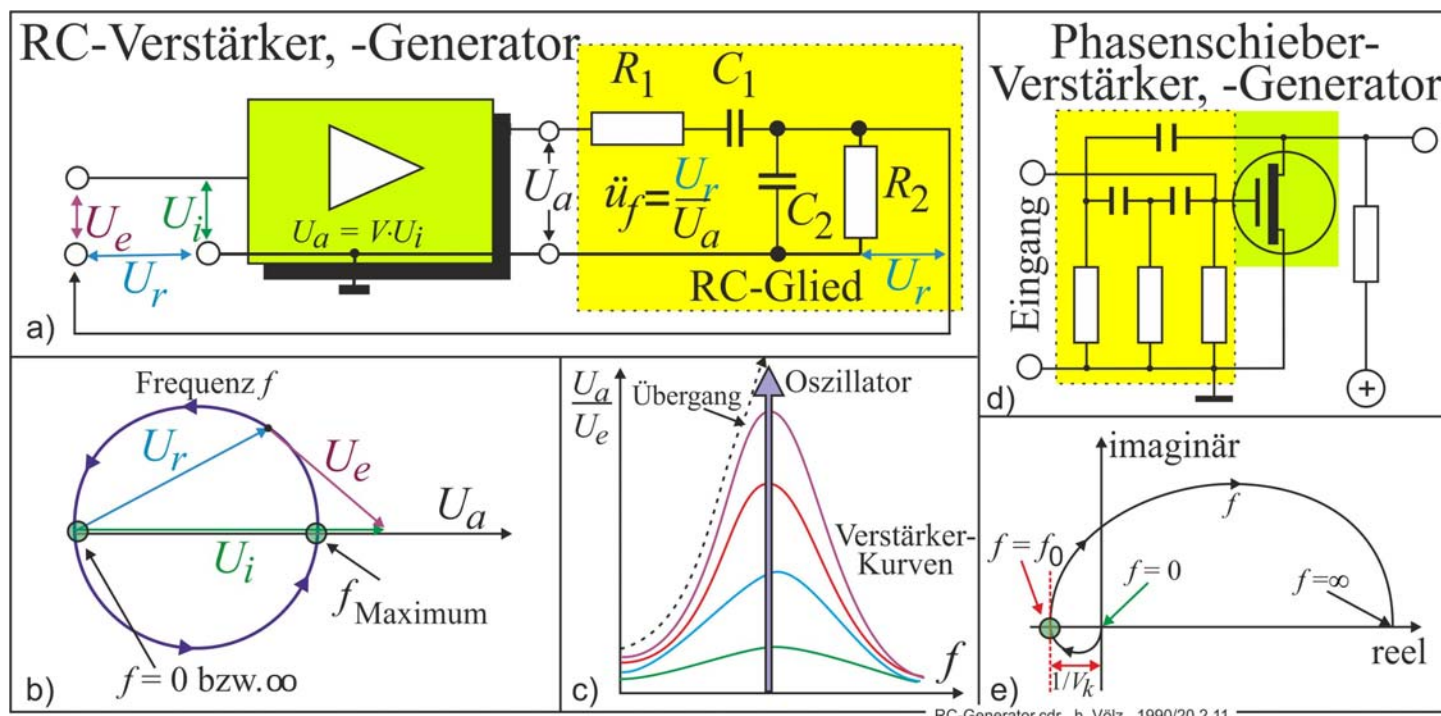


Bild 38. Zwei Beispiele für selektive RC-Verstärker und RC-Generatoren.

Bei der RC-Variante (a) müssen U_i und U_a gleichphasig sein. Das verlangt einen zweistufigen Verstärker. Für die Variante (d) genügt dagegen ein Transistor. Dafür muss aber das RC-Glied für die maximale Frequenz f_0 die Phase um 180° drehen. Im Mittel muss dann jeder einzelne Hochpass etwa 60° realisieren. Die Ortskurve des dreistufigen Hochpasses zeigt (e). Bei f_0 dämpft jedes einzelne Glied die Amplitude etwa auf $1/2$, alle zusammen etwa auf $1/8$. Für einen RC-Generator muss daher die Verstärkung des Transistors $V \geq 8$ betragen. Das ist erheblich mehr als im Fall (a). Bei dieser Schaltung wird aber noch deutlicher, dass Resonanzeffekte und Schwingungen durchaus nicht **L** und **C** oder auch nur Frequenzmaxima benötigen. Entscheidend ist die selektiv phasenrichtige Rückkopplung mit 0 oder 180° .

²¹ Lateinisch generator Erzeuger.

²² Lateinisch oscillare sich schaukeln, oscillum Schaukel.

In den meisten Fällen werden Sinusschwingungen mit Hilfe von **Schwingkreisen** erzeugt. Die dabei wichtigsten Fälle zeigt **Bild 39**. Zeitlich nacheinander entstanden dabei mehrere Schaltungen, welche die Namen der Entwickler tragen. Bei allen treten ähnlich wie bei den RC-Generatoren auch Schwankungen um die Resonanzfrequenz auf, die jedoch deutlich geringer sind. Mittels der mechanischen Resonanz eines Quarzes können sie noch erheblich weiter vermindert werden. U. a. bei den Armbanduhren wird der Quarz sogar zu einer minimalen Stimmgabel geformt.

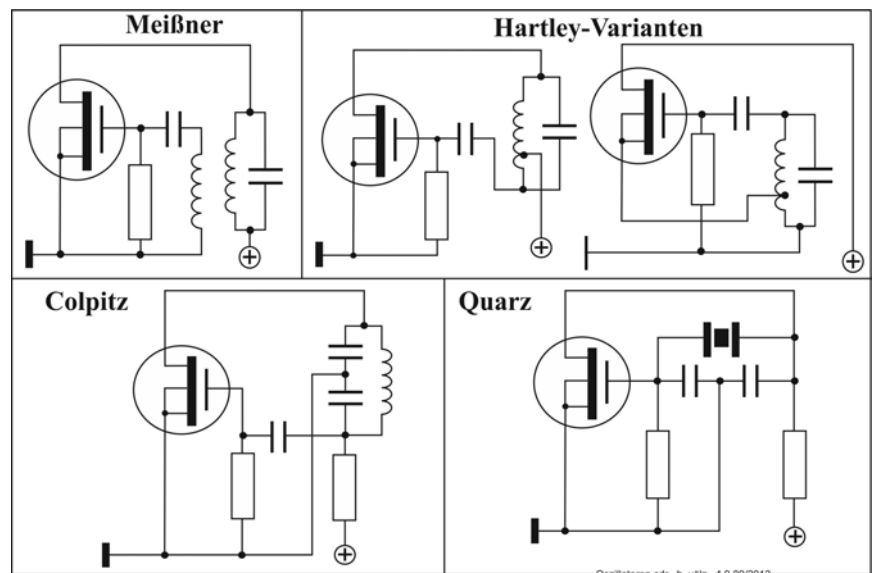


Bild 39. Varianten der Oszillatoren mit Schwingkreis.

Periodische, aber *nicht sinusförmige Schwingungen* treten in anderen Fällen auf. Ein typisches Beispiel ist der **Multivibrator**, der auch Rechteckgenerator genannt wird. Leider ist sein Geschehen schwierig zu beschreiben und fast unmöglich zu erklären. Es sei hier mit dem **Bild 40** versucht. Die beiden Transistoren sind gegenseitig über die beiden Kondensatoren extrem stark rückgekoppelt. Beim Einschalten der Betriebsspannung sei durch Unsymmetrie oder Zufall der Strom durch den Transistor T_1 größer als bei T_2 . Das bewirkt ein stärkeres Sinken der Spannung an seinem Kollektor mit der Folge, dass T_2 einen Impuls in Sperrrichtung erhält. Was an seinem Kollektor eine Spannungserhöhung bewirkt. Über den Kondensator wird dadurch T_1 zu weiter erhöhten Strom angesteuert und so immer weiter fort, bis schließlich keine Stromerhöhung mehr möglich ist. So kann T_2 über den Widerstand an seiner Basis stärker Strom ziehen. Dadurch kehrt sich der soeben beschriebene Vorgang um, usw. auch in Gegenrichtung. Als Folge entsteht die typische periodische Rechteckschwingung. Die Spannungen an den verschiedenen Stellen sind dabei in den Kreisen als Signalverläufe dargestellt.

Werden die beiden Kondensatoren jedoch durch Widerstände ersetzt, so entsteht ein **FlipFlop**. Durch die o. g. Unsymmetrie oder den Zufall wird dann beim Einschalten der Betriebsspannung wieder einer der beiden Transistoren stärker leitend und damit der andere verstärkt in Sperrrichtung getrieben. Doch dieser Zustand bleibt dann erhalten. Er kann aber bei Bedarf durch einen externen Impuls umgekehrt werden. Stabil bestehen nur die beiden extrem entgegengesetzten Zustände. Das zeigt schematisch (b). Liegt beim Einschalten ein Wert oberhalb der Trennlinie vor, so wird immer T_1 leitend werden, bei unterhalb aber T_2 . Daher speichert diese Schaltung ein Bit. Diese FlipFlop können gemäß (c) auch durch zwei rückgekoppelte Logikschaltungen erzeugt werden.

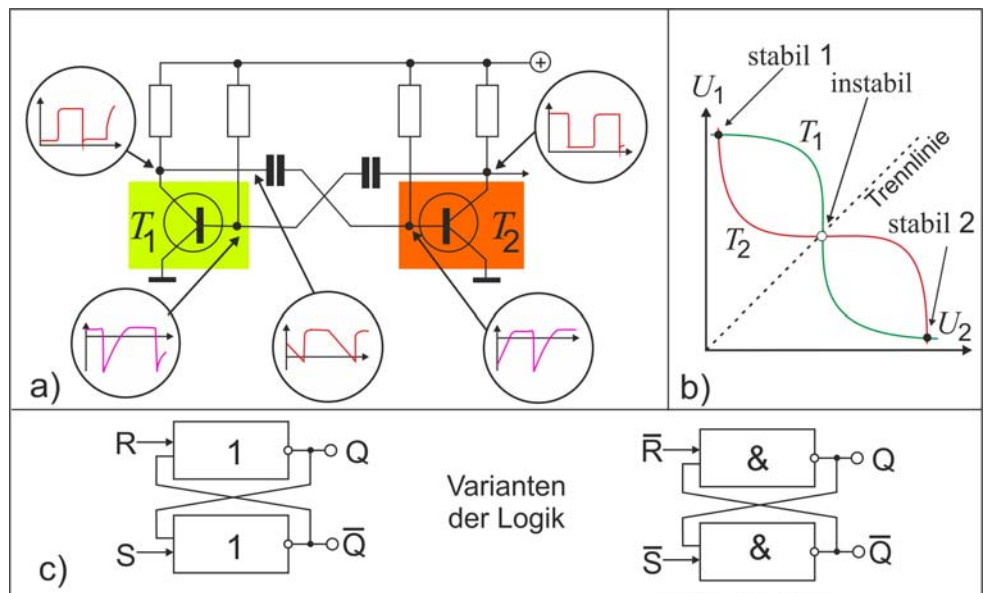
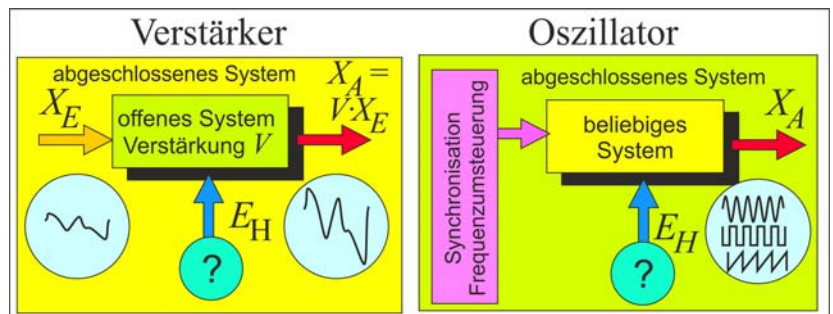


Bild 40. Zur Erklärung von Multivibratoren und FlipFlop.

Im Kontext der Verstärker wurden bisher nur ausgewählte elektronische Oszillatoren beschrieben. Doch bezüglich der Röhren und Transistoren besteht zwischen beiden ein recht allgemeiner Zusammenhang, den **Bild 41** hervorhebt. Beide benötigen meist – sofern sie nicht intern vorhanden ist – eine Hilfsenergie E_H zum Betrieb. Ein Verstärker verlangt immer ein Eingangssignal, damit er es vergrößert am Ausgang bereitstellt. Beim Oszillator ist es generell nicht notwendig. Spezielle Eingangssignale können jedoch seine Frequenz auf eine externe synchronisieren oder die Kurvenform seines Ausgangssignals verändern.

Bild 41. Vergleich von Verstärker und Oszillator,



Stabilisierung

Bei vielen Anwendungen ist es vorteilhaft bis notwendig bestimmte Parameter, z. B. die Temperatur konstant zu halten. Am einfachsten leistet das eine Abschirmung der Störeinflüsse. Doch das ist nicht immer möglich. Bereits auf S. 18 wurde stattdessen bereits auf die schon im Altertum bekannte Regelung bei der Öllampe und für eine konstante Durchflussmenge mit dem Schwimmventil hingewiesen. Als noch eine einzige Dampfmaschine für den Antrieb vieler Verarbeitungsmaschinen über eine Transmissionswelle (**Bild 42**) erfolgte, war es schwierig die Drehzahl beim An- oder Abschalten einzelner Maschinen konstant zu halten. Hierfür erfand 1769 James Watt seinen Fliehkraftregler (**Bild 43**). Mit steigender Drehzahl werden die rotierenden Kugeln als Folge erhöhter Fliehkraft weiter nach außen getrieben und reduzieren so mittelbar den Dampfzustrom. Bei sinkender Drehzahl fallen die Kugeln etwas zurück und erhöhen dann die Dampfzufuhr. Für die Elektronik sind oft gut stabilisierte Spannungen oder

Ströme erforderlich. Das leisten die Reglerschaltungen von **Bild 44**. In (a) wird aus der schwankenden, aber zu großen Eingangsspannung U_e mittels einer Z-Diode eine konstante Bezugsspannung hergestellt, die jedoch nicht belastbar ist. Deshalb wird nach dem Regler mit dem Spannungsteiler aus U_a eine Vergleichsspannung hergestellt. Beide Spannungen werden dem OV zugeführt, der je nach ihrer Differenz den Regler so verstellt, dass unabhängig von der Belastung und den Änderungen der Eingangsspannung U_e immer genau die gewünschte Ausgangsspannung U_a bereit steht. Zur Sicherheit vor Überlastung können noch Zusatzschaltungen eine Strombegrenzung erforderlich machen (c). Bei der Konstantstromschaltung (b) erfolgt die Regelung durch den Vergleich mit dem Spannungsabfall an R im Stromzweig von I_a .

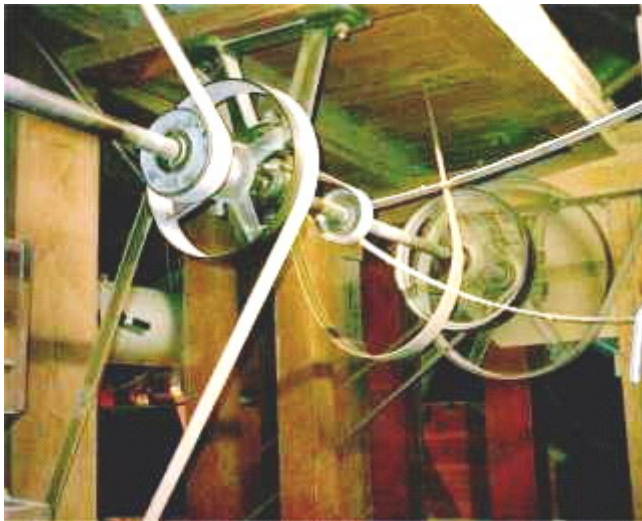


Bild 42 Bildausschnitt einer typischen Transmissionswelle (links).

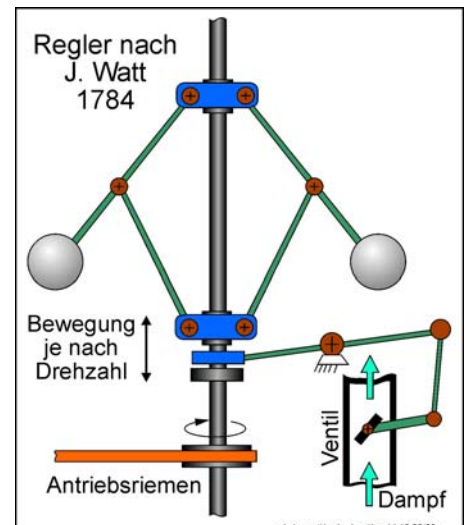


Bild 43. Prinzip des Drehzahlreglers nach Watt (rechts).

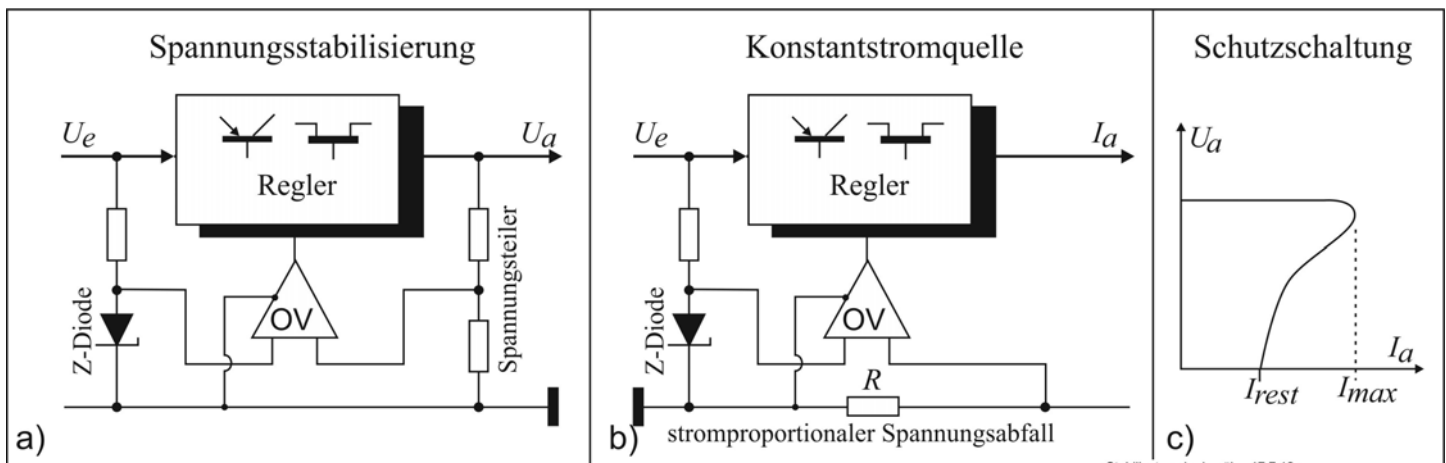


Bild 44. Konstantquellen für Spannung (a) und Strom (b) sowie Kennlinie einer Schutzschaltung für zu hohe Last (c).

In der Kybernetik existieren zur Konstanthaltung beliebiger Werte die **Steuerung und Regelung**. Am Beispiel der Zimmertemperatur zeigt **Bild 45** ihren Unterschied: Bei der Steuerung (a) werden möglichst viele Einflusswerte erfasst und daraus dann das Ein- bzw. Ausschalten der Raumheizung bewirkt. Da es immer zusätzliche Einflussgrößen gibt, funktioniert das nur ungefähr. Meist arbeitet die Regelung (b) genauer. Sobald der Anzeigewert des Thermometers den festgelegten Schaltpunkt unterschreitet wird die Heizung eingeschaltet und beim Überschreiten ausgeschaltet. Das führt immer zu einer gewissen Verzögerung und zu einer etwas unsicheren Temperatur. Das ist recht gut beim Kühlschrank bekannt. Die üblich eingestellte Temperatur schwankt etwa zwischen 9 und 11 °C. Deshalb entstand die Regelung mit **Störwertaufschaltung**. Für die Beeinflussung von Stoff-, Energie- und Informationsströmen wird sie mit den einfachen Regelungen und Steuerungen in **Bild 46** verglichen. Dem Regelverstärker werden dabei zusätzlich Messwerte der Störungen und Schwankungen der Eingangsgröße passgerecht zugeführt. Dadurch wird der Toleranzbereich deutlich verringert. [Völ66].

Bild 45. Regelung und Steuerung der Zimmertemperatur (rechts oben).

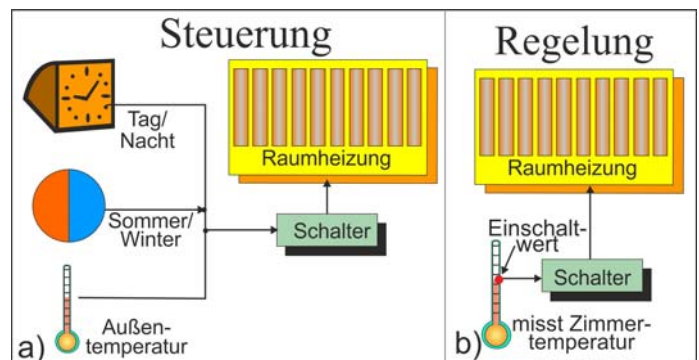
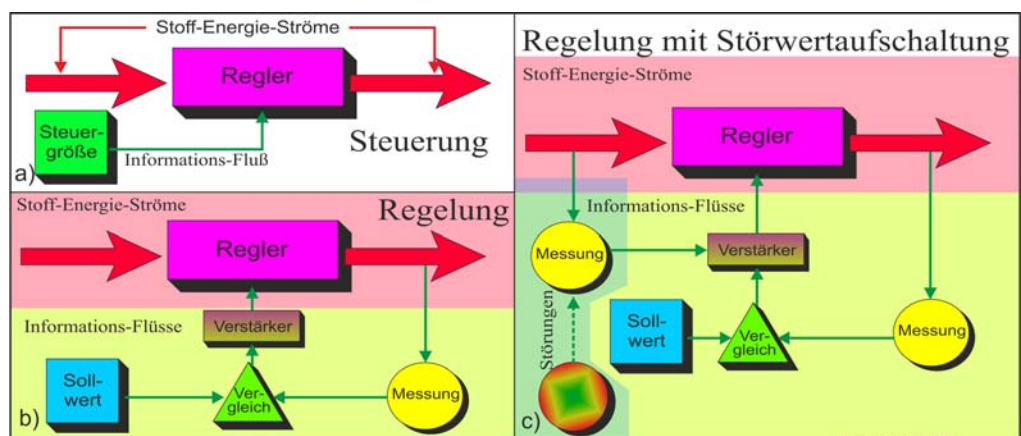


Bild 46. Vergleich der Varianten zur Konstanthaltung einer Größe.



Tunneldiode und Gunn-Effekt

Bei der Rückkopplung auf S. 19 und mit Bild 33b ist gezeigt, wie negative Widerstände entstehen und was sie Besonderes zu leisten vermögen. Sie kommen auch noch bei Gasentladung (Bild 10, S. 7) und bei zwei besonders vorteilhaften, speziellen Halbleiterbauelementen vor: Der Tunneldiode und dem Gunn-Element. Leider gibt es für negative Widerstände keine brauchbare **mechanische Analogie**. Bestenfalls können noch Ölen, Schmieren usw. zur Vermeidung von Reibung so betrachtet werden.

Die erste **Tunneldiode** (auch Esaki- oder backward-Diode genannt) wurde 1958 von Leo Esaki gebaut. Das dafür notwendige quantenphysikalische Tunneln hatte bereits 1954 William Bradford Shockley vermutet. Gegenüber anderen Dioden besitzt sie erstens einen extrem abrupten pn-Übergang, der in weniger als $0,01\ \mu\text{m}$ erfolgt und zweitens eine hohe Dotierung von $>2,5 \cdot 10^{19}$ Störstellen/ cm^3 . Das Valenz- und Leitungsband sind so durch eine sehr schmale Sperrzone getrennt. So können Ladungsträger den Potentialwall der Sperrzone quantenmechanisch ohne Energieverlust durchtunneln (Bild 47e) und relativ große Ströme hervorrufen. Das führt zum steil linearen Anstieg der Strom-Spannungs-Kennlinie aus dem Ursprung gemäß (a). Die abnehmende Überlappung der beiden Bänder führt dann dazu, dass der Tunnelstrom nach seinem Maximum auf Null absinkt (blau). Wird die Tunneldiode weiter in Durchlassrichtung betrieben, so tritt der übliche Durchlassstrom einer Diode auf (rot). Insgesamt entsteht so ein Kennlinienteil, in dem mit steigender Spannung der Strom abnimmt. Das bewirkt ihren differentiellen negativen Widerstand zwischen U_{max} und U_{min} . Meist liegt die minimale Spannung bei etwa 50 mV und die maximale zwischen 300 und 500 mV. Das Verhältnis $I_{\text{max}}/I_{\text{min}}$ ist materialabhängig: Si ≈ 6 , Ge ≈ 10 und GaAs ≈ 60 (d). Mit der Größe der Diode kann ein maximaler Strom von $<1\text{mA}$ bis zu einigen Ampere erreicht werden. Da der Tunneleffekt in $<1\text{ ns}$ erfolgt, ist die Kennlinie bis in den GHz-Bereich voll gültig. Bei den meisten Anwendungen wird ihr negativer Widerstand zur Entdämpfung und Schwingungserzeugung im Sinne von Bild 33c, S. 19 benutzt. Die zweite wichtige Anwendung erfolgt in der Digitaltechnik. Sie wird dann über einen Vorwiderstand R_V aus einer Gleichspannung U_B betrieben (c). Die Widerstandsgerade schneidet dann unter bestimmten Bedingungen dreimal die Kennlinie der Tunneldiode (b). Lediglich die Arbeitspunkte 1 und 3 sind stabil und bedingen logisch 0 oder 1.

Der **Gunn-Effekt** wurde 1963 von John Battiscombe Gunn entdeckt. Er tritt vor allem in Galliumarsenid (GaAs) oder Indiumphosphid (InP) bei hohen elektrischen Feldstärken auf und besitzt stückweise einen negativen Widerstand. (Bild 47 f, g). Das Gunn-System enthält keinen pn-Übergang. Deshalb ist die übliche Bezeichnung Gunn-Diode eigentlich falsch und wäre besser, aber umständlicher als Elektronentransfer- oder Mikrowellenzweipol-Element beschrieben. Bei Feldstärken $>2\text{ kV/cm}$ treten im Material Mikrowellen-Schwingungen auf. Sie werden durch die lokal gekrümmten Bandstrukturen hervorgerufen (a), indem sie Ladungsträger mit unterschiedlicher Beweglichkeit entstehen lassen (b). Bei entsprechender Beschleunigung verdichten sich die Ladungsträger zu Ladungswolken, die durch das Material wandern und so die Mikrowellen bewirken. Der differentiell negative Ast betrifft hier also nicht den Widerstand, sondern die Geschwindigkeit. In den 60er und 70er Jahren wurde versucht, diesen Effekt auch zur Speicherung zu nutzen. Eine praktische Anwendung ist jedoch nicht entstanden.

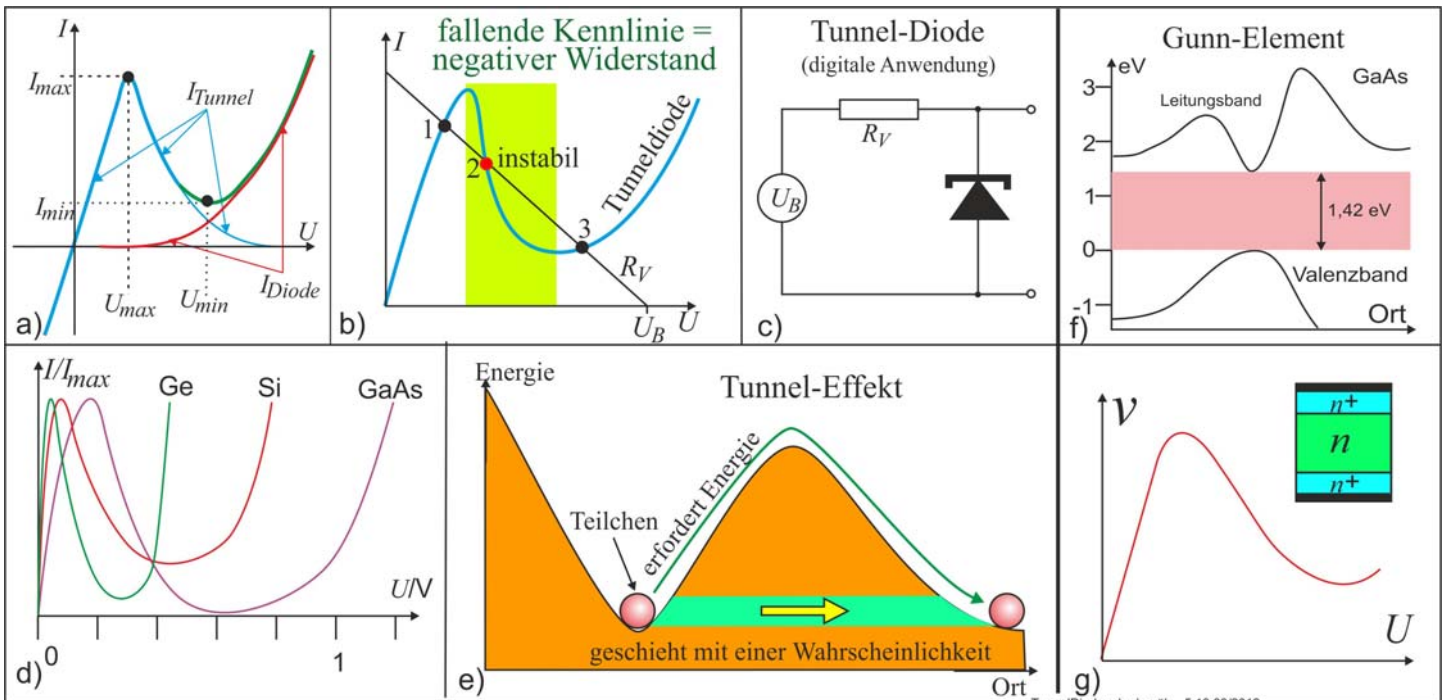


Bild 47. Zu den Eigenschaften von Tunneldioden und Gunn-Bauelementen

5. E-Verstärker durch Effekte

Maser und Laser

Beide Begriffe sind Abkürzungen: **microwave bzw. light amplification by stimulated emission of radiation**, also etwa Mikrowellen- bzw. Lichtverstärkung durch stimulierte Emission von Strahlung. Die stimulierte Emission hatte Einstein 1917 vorausgesagt. 1928 haben sie dann Rudolf Ladenburg und Hans Koppermann bei Gasentladungslampen nachgewiesen. Maser und Laser unterscheiden sich vor allem durch die typischen Frequenzen: Den Maser gibt es für etwa 100 kHz bis 100 GHz; den Laser im sichtbaren Bereich von 400 bis 800 nm, aber auch für Infrarot und UV, sogar für Röntgenstrahlung. Beim Maser sind auch die Begriffe wie **Molekular- oder Quanten-Verstärker** gebräuchlich. Ihre **Anwendungen** sind vielfältig u. a. Werkstoffbearbeitung, Medizin, Messtechnik, Wissenschaft, Holographie, Daten- und Militärtechnik, spezieller Lasershow, RGB-Projektoren und Laserpointer. Im Folgenden wird bevorzugt Laser auch dann benutzt, wenn es infolge der Wellenlänge eigentlich bereits Maser sind.

Entsprechend der Quantentheorie besteht Licht aus **Photonen** (Lichtquanten) mit der Energie $E = h \cdot \nu$; dabei sind ν die Lichtfrequenz und h die Planck-Konstante. Klassisch betrachtet ist ein Photon eine elektromagnetische **Transversal-Welle** endlicher Länge, die durch eine Schwingungsrichtung gekennzeichnet ist. Mehrere Photonen können polarisiert in einer Richtung schwingen oder unpolarisiert sein (Bild 48a). Weißes Licht besitzt Schwingungen aller Frequenzen im sichtbaren Bereich. Bei **monochromatischem** Licht besitzen alle Photonen die gleiche Frequenz, können sich aber an verschiedenen Orten befinden (b). Mit dem Laser oder Maser können mehrere Photonen (gleicher Richtung) phasenstarr aneinander gekoppelt werden. Dabei ergibt sich die **Kohärenzlänge** (c), innerhalb der Interferenzen auftreten können.

Licht kann quantenphysikalisch **absorbiert oder abgestrahlt** werden. In (d) ist der Grundzustand eines Atoms als grüne Linie und der angeregte als rote Linie angedeutet. Bei der Absorption hebt das Photon (grüne Wellenlinie) ein Elektron auf das höhere Niveau. Bei der spontanen Lichtemission sinkt es in den Grundzustand zurück und strahlt dabei ein Photon ab. Bei der stimulierten Emission befindet sich ein Elektron auf dem höheren Niveau. Ein passgerecht eingestrahktes Photon bewirkt dann, dass das Elektron in den Grundzustand gelangt und dabei ein zweites Photon zum eingestrahkten kohärent hinzufügt.

Im quantenphysikalischen Bändermodell ist der Abstand zwischen den Energietermen gekennzeichnet durch die Wellenzahl ortsabhängig (g). Entsprechend (h) können dabei das Minimum des höheren gegenüber dem Maximum des unteren Terms seitlich verschoben sein (h). Dann benötigt das Elektron zum Herunterfallen einen seitlichen Impuls, der vor allem thermisch entsteht und daher zu erheblichen Verzögerungen führen kann. Es liegt dann ein **indirekter Übergang** vor.

Voraussetzung für die Funktion des Lasers ist ein 3-Term-System (i). Unter „normalen“ Bedingungen stellt sich dann eine thermische Verteilung der Belegungen ein (i), die der thermodynamischen Boltzmann-Verteilung entspricht. Besitzt dagegen E_2 einen indirekten und damit metastabilen Übergang, so wird seine Belegung deutlich größer werden. Denn die meisten Elektronen die durch Pumpen (Energiezufuhr) nach E_3 gelangen, fallen auf E_2 zurück und können diesen Term kaum verlassen (j). Eine metastabile Belegung kann unterschiedlich lange, bis zu Stunden dauern. Es ist als Nachleuchte bei Lumineszenz, Phosphoreszenz usw. bekannt. Der Strahlungsübergang zu E_1 kann aber mittels eines stimulierenden Signalstrahls ausgelöst werden. Damit dann schlagartig alle vorhandenen Elektronen folgen, ist ein Interferometer²³ (oder Hohlraumresonator besonders bei Maser) notwendig (e). Es ist beidseitig mit teildurchlässigen Spiegeln versehen, Zwischen ihnen bildet sich beim ersten herunterfallenden Elektron eine stehende Welle aus, die dann schlagartig über stimulierte Emission alle Elektronen auf E_2 nachzieht. So entsteht ein intensiver Laserimpuls. Für einen absichtlichen Start kann über den linken teildurchlässigen Spiegel ein Signalstrahl die erste stimulierte Emission auslösen. Der danach rechts austretende Laserimpuls ist das erheblich verstärkte Signal (e). Ist der linke Spiegel undurchlässig, so ist kein Signal zum Auslösen möglich. Dann bewirkt die thermische Energie rein zufällig den Laserimpuls (h). Durch fortlaufendes intensives Pumpen (als Hilfsenergie) entsteht ein anhaltender Laserstrahl. Der Laser ist zum Oszillator für die typische Lichtfrequenz geworden. Im Interferometer ist stets eine deutlich höher Energiedichte als beim austretenden Strahl vorhanden. Das bedeutet eine hohe Materialbelastung und verringerte Lebensdauer.

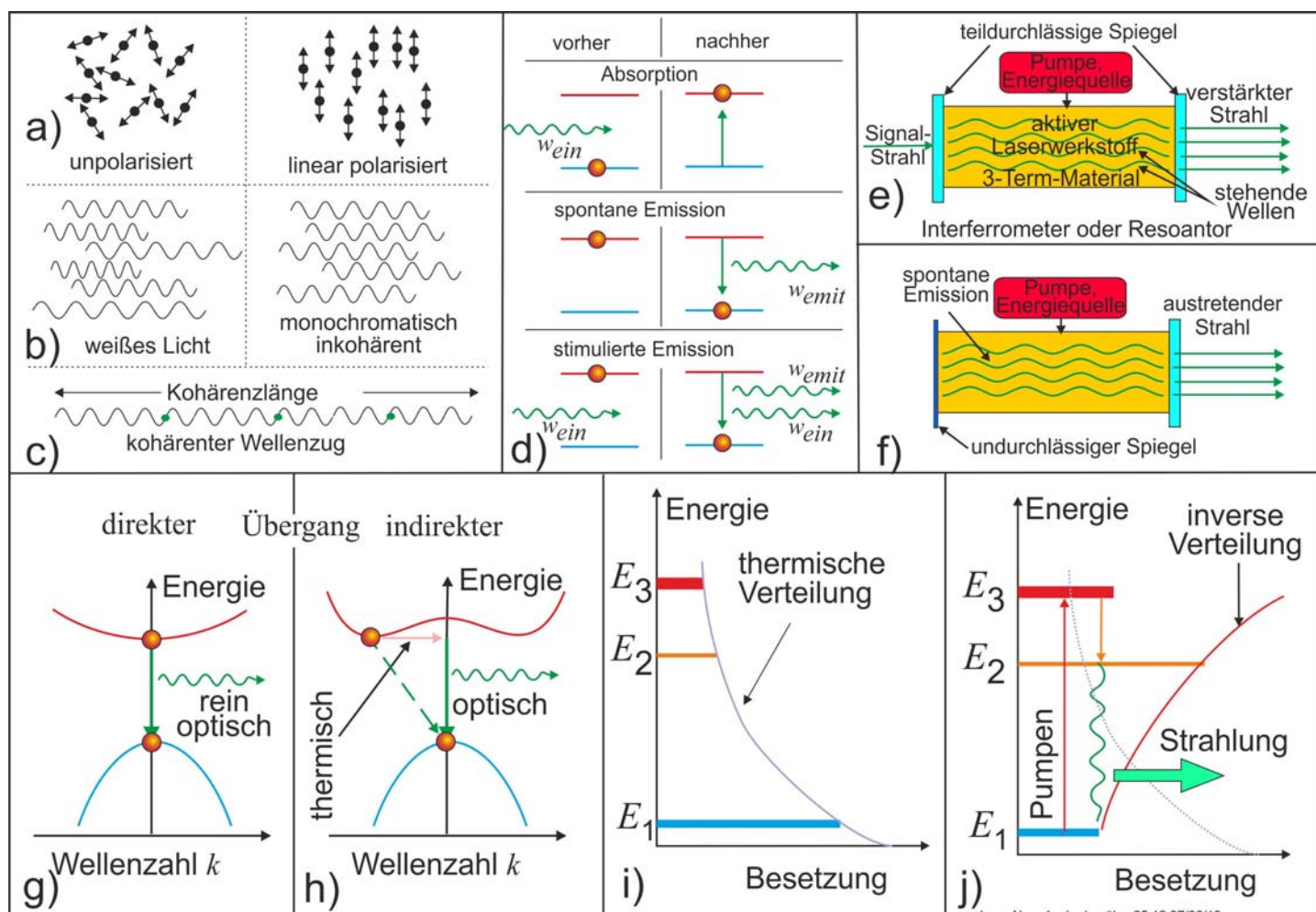


Bild 48. Zur Erklärung des Lasers

Es sind mehrere **Laser-Varianten** entstanden, die sich vor allen durch das Lasermaterial und den Pumpvorhang unterscheiden (f_p ist die Pumpfrequenz). die wichtigsten Klassen sind:

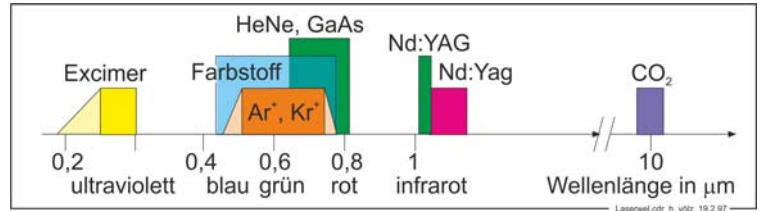
- **Halbleiter- = Injektions-Laser**, besitzen negativen Temperaturkoeffizienten, daher Gefahr der Selbsterstörung, Strombegrenzung notwendig. Problem bei Intensitätssteuerung, oft nur ein- und ausschaltbar; **Pumpen**: Stromfluss eines *pn-Übergangs*
- **Festkörper-Laser**: speziell dotierte Kristalle, z. B. **Rubin** durch geringe Mengen Cr_2O_3 tiefrot gefärbt ($\alpha\text{-Al}_2\text{O}_3$ 694,3 nm), $f_p = 30$ GHz, Leistungsverstärkung 10^3 , Bandbreite 50 MHz, Magnetgleichfeld ≈ 3 kOe, magnetisch durchstimmbar; **YAG** (Yttrium-Aluminium-Granat (≈ 350 MHz), **Neodym** (1,064 μm), **Ytterbium** (940 nm) und **Titan** (670 - 1100 nm). **Pumpen** mit energiereichen Blitzröhren.
- **Gas-Laser**: Resonator mit Gas gefüllt ($p = 10$ bis 10^6 Pa): **HeNe** wichtiger Industrie-Laser (632 nm), **CO₂** (11 μm), **CO** (6–8 μm), **N₂** (337 nm), **Ar** (mehrere Linien 450 – 500 nm), **He** (442 + 325 nm), **NH₃** (12,7 mm, Maser). **Pumpen** durch elektrische Gasentladung, seltener Mikrowellen.

²³ Interferometer: lateinisch ferire schlagen, treffen. Resonator: lateinisch re- wieder, zurück; sonare tönen, hallen; resonare wieder ertönen.

- **Farbstoff-Laser** nutzen organische Farbstoffe in alkoholischer Lösung (oft Methanol oder Ethanol) oder durchsichtigen Werkstoff mit Farbzentren. zuweilen durchstimmbar. **Pumpen** mit Blitzröhren oder Stickstoff-Gas-Laser.
- **Chemische Laser** nutzen zum **Pumpen** chemische Reaktionen.
- **Maser** nutzen u. a. Molekülschwingungen oder magnetische Dipolübergängen in Atomen, Vorwiegend paramagnetische Ionen, u. a. Cr-, Fe-, Gd- und Ni-Ionen. Kristall in Hohlraumresonator, oft tiefe Temperatur flüssiges He (wenige K).

Einen Überblick zum Zusammenhang zwischen dem Lasermaterial und den verstärkten bzw. erzeugten Frequenzen zeigt **Bild 49**.

Bild 49. Material und Frequenzbereiche bei Lasern.



Auch beim Laser und Maser ist der Übergang vom **Verstärker zum Oszillator** gut zu erkennen. (Bild 64e, f). Dabei wird sogar die hochselektive Verstärkung des thermischen Rauschens (s. S. 21 bei Rückkopplung) noch deutlicher und gut belegbar (**Bild 50**). Da die „Länge“ des Resonators meist sehr groß gegenüber der erzeugten Wellenlänge ist, kann sich jedes Mal beim Beginn der stimulierten Emission spontan eine etwas andere Anzahl von Wellenbergen für die stehende Welle ergeben. Mit größerer Intensität (analog höher Rückkopplung) sinkt daher deutlich ihre Anzahl und damit im Mittel die Bandbreite der Schwingung (b und c). Dazu ergänzt (a) den Vergleich des Anstieges der Strahlungsintensität bei der LED und dem Laser. Ergänzend ist noch (d bis f) gezeigt, wie durch Änderung und Weiterentwicklung des Resonators und des pn-Übergangs die Richtung und der Öffnungswinkel der Strahlung vorteilhaft verändert wurden.

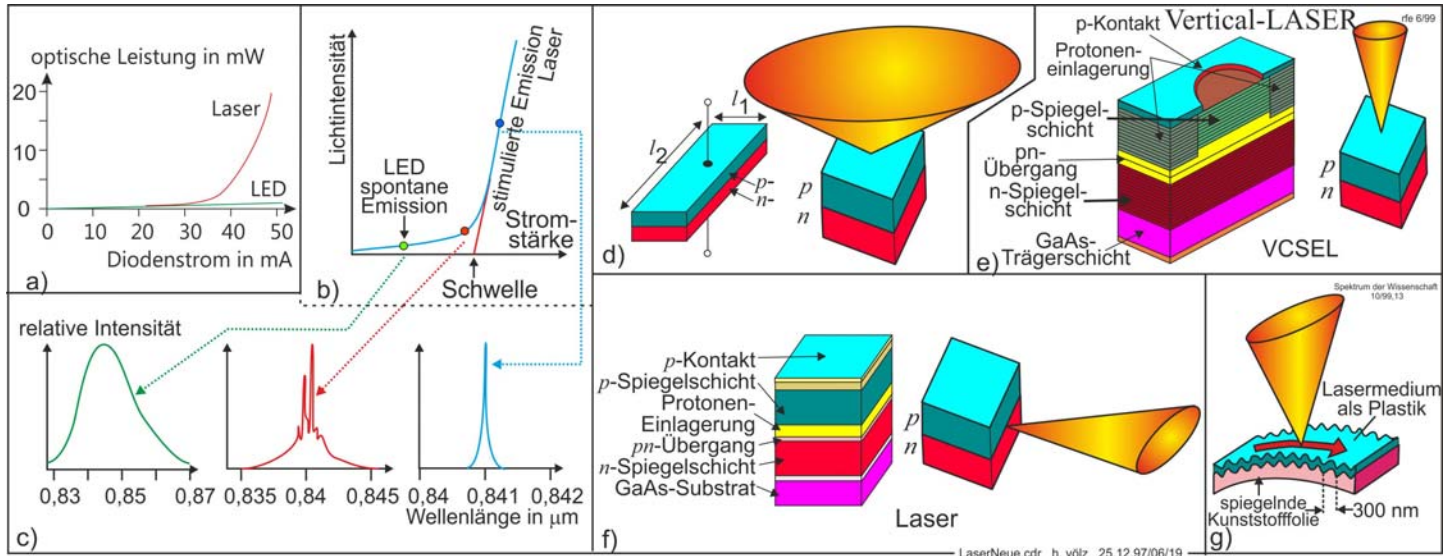


Bild 50. Zum Spektrum und dem Abstrahlungskegel der Laserstrahlung.

Aktive Bauelemente für Mikrowellen

Die üblichen Röhren und Transistoren sind nur bis zu einer **oberen Frequenzgrenze** im MHz-Bereich nutzbar. Höhere Frequenzen sind wegen der **Laufzeit der Elektronen** von der Kathode zur Anode bzw. von Source nach Drain nicht erreichbar. Selbst extremes Reduzieren der Abstände wie bei der Scheibentriode (Bild 16, S. 10) und bei speziellen Kurzkanal-Transistoren ermöglicht nur eine geringe Steigerung der oberen Grenzfrequenz. zusätzlich wirken sich auch unvermeidliche Kapazitäten und Induktivitäten aus. Daher entstanden mehrere Varianten, welche die Laufzeit direkt nutzen. Ihre wichtigsten Varianten sind in **Bild 51a** systematisch erfasst. Vor allem sind das Trift- (Klystrons) und Wanderfeldröhren, u. a. als Magnetrons. Im anderen Kontext wurden bereits die Gunn- und Tunnel diode (Bild 47, S. 24) sowie die Scheibentriode (s. o.) behandelt. Eine sehr spezielle Variante zur Triode, die heute nicht mehr gebräuchlich ist, betrifft den Barkhausen-Kurz-Effekt. Bei ihr wird an das Gitter eine positive und an die Anode eine negative Spannung angelegt. Dadurch fliegen die Elektronen beschleunigt durch das Gitter hindurch zur Anode und werden von dort zum Gitter reflektiert. Beschleunigt fliegen sie durch deren Lücken, werden dann von der Katode erneut reflektiert usw. Dabei entstehen sehr hohe Frequenzen, die durch externe Resonatoren festgelegt und verstärkt werden. Bei den meisten Mikrowellen-Röhren besteht eine gewisse Analogie zum Laser/Maser (b). Meist wirken dabei Resonatoren verstärkend oder ermöglichen die Selbsterregung.

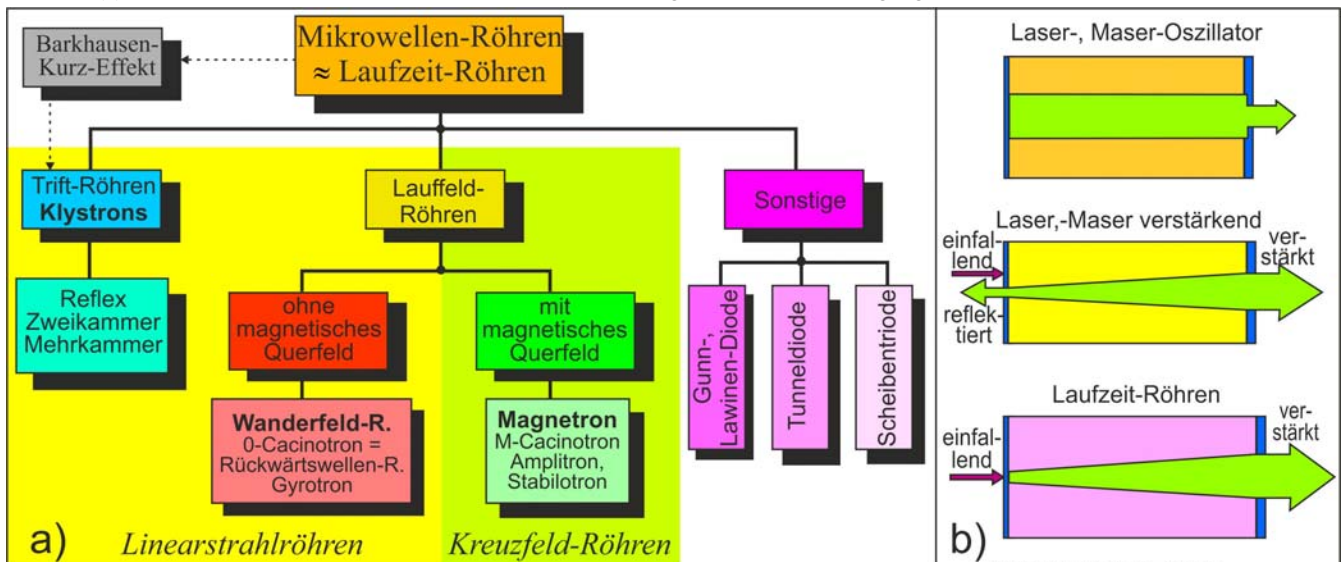


Bild 51. Überblick zu den Mikrowellen-Röhren.

Laufzeiteffekte sind besonders deutlich beim **Zwei-Kammer-Klystron**²⁴ (Triftröhre) zu erkennen (**Bild 52a**). Es besteht im Wesentlichen aus einer Elektronen-Strahl-Kanone und zwei Hohlraum-Resonatoren. Die Elektronen durchfliegen nacheinander die beiden Resonatoren. Das aktuelle Feld im ersten Resonator beschleunigt oder bremst ihre Geschwindigkeit. Während der Laufzeit (Driftstrecke) überholen dabei die schnellen Elektronen die langsamen und es entsteht eine **Dichtemodulation** des Elektronenstrahls (b). Im Resonanzfall treten die Elektronen-Pakete (= bunche) besonders ausgeprägt beim zweiten Resonator auf und er wird zum verstärkten Schwingungen angeregt. Hier werden sie ausgekoppelt, dieses **Lawinen-, Laufzeit-Prinzip** wurde 1958 von W. T. READ gefunden.

Mehrkammerklystrons ermöglichen besonders hohe Verstärkerleistungen und größere Bandbreite. Dabei werden die hintereinander angeordneten Resonatoren als unterschiedlich abgestimmte Triftstrecken gesteuert.

Ein **Reflexklystron** (c) benötigt nur einen Hohlraum-Resonator. Durch eine Anode mit negativer Spannung werden die Elektronen ähnlich wie beim Barkhausen -Kurz-Effekt (s. o.) reflektiert. Auch dabei entstehen Elektronen-Pakte, die nun aber erneut, jedoch verstärkt den einzigen Resonator erreichen und so entsteht Selbsterregung. Das Reflexklystron ist daher ein Oszillator für hohe Frequenzen. Mittels Änderung der Reflektorspannung ist unmittelbar Frequenzmodulation möglich. Oft ist zusätzlich der Hohlraum auch mechanisch verstimmbar.

der dichtemodulierte Elektronenstrahl eines Klystrons enthält viele Oberwellen, Das ermöglicht leistungsfähige **Frequenzvervielfacher**. Klystrons ermöglichen Frequenzen von 0,5 bis 500 GHz. Ihre Dauerstrich-Leistungen reichen bis 1 MW, Impulsleistungen bis 100 MW.

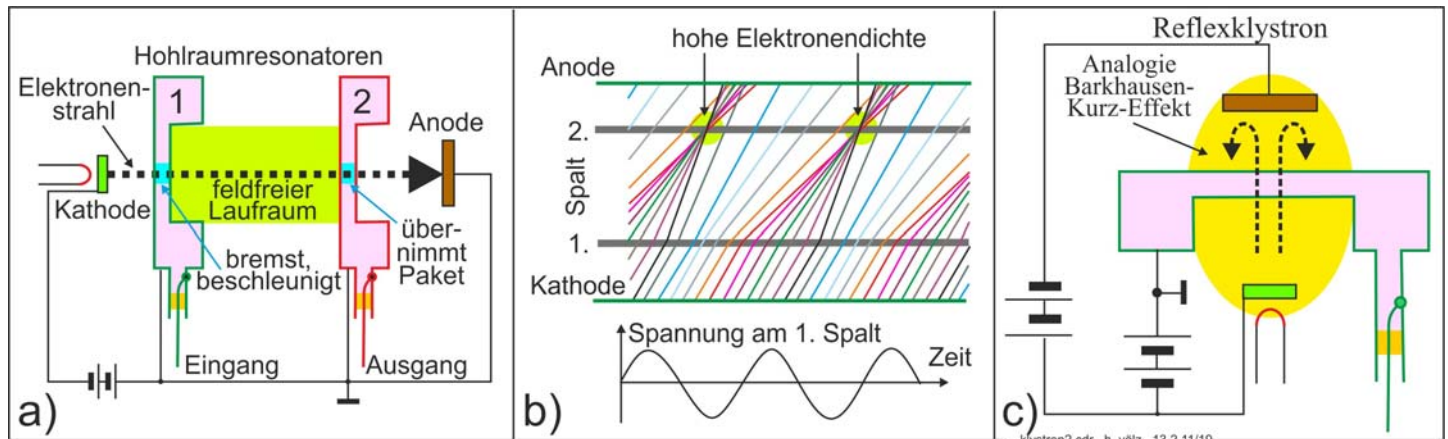


Bild 52. Aufbau eines Klystrons (a), die erfolgende Dichtemodulation (b) und Schema eines Reflexklystrons (c).

Die Wanderfeld-Röhre (Travelling Wave Tube, TWT) nutzt die Wechselwirkung zwischen Elektronenstrahl und einem bewegtem elektrischem Feld (**Bild 53**). Ein elektrisches Feld breitet sich etwa mit Lichtgeschwindigkeit ($\approx 3 \cdot 10^8$ m/s) aus. Elektronen fliegen dagegen erheblich langsamer. Bei 1 kV Beschleunigung werden z. B. nur $v_e \approx 7 \cdot 10^6$ m/s erreicht. Bei der TWT wird die Feldausbreitung mittels eines spiralförmigen Leiters (Helix) auf v_h reduziert, die dennoch nur wenig kleiner als v_e ist. Der Elektronenstrahl wird mit einem statischen Magnetfeld von $\approx 0,1$ Tesla gut fokussiert durch Helix geleitet. Die bewegten Elektronen übertragen dabei Energie auf die Helix und verstärken so ein dort vorhandenes Signal. Wegen $v_e < v_h$ tritt der Cerenkov-Effekt auf und infolge der Laufzeit werden ähnlich wie Klystron Pakete (bunche) erzeugt. Die Ein- und Auskopplung des Signals erfolgt mittels koaxialen Anschluss am den Hohlraum-Resonatoren oder Hohlleitern. TWT erzeugen im Bereich von 1 bis 20 GHz nur kleine bis mittlere Ausgangsleistungen. Andererseits arbeiten sie trotz großer Verstärkung (>20 dB) sehr breitbandig ≈ 300 MHz. Die hohe Verstärkung kann leicht zur Selbsterregung führen, nämlich dann, wenn auf der Wendel zurücklaufende Signale wirksam werden. Um das zu vermeiden wird oft in der Röhrenmitte ein Dämpfungsglied eingefügt. Verwandte Typen – die aber hier nicht näher besprochen werden – sind u. a. das **Carcinotron** (Rückwärtswellenröhre) als Oszillator und das **Gyrottron** als Generator bis herab zu Millimeterwellen.

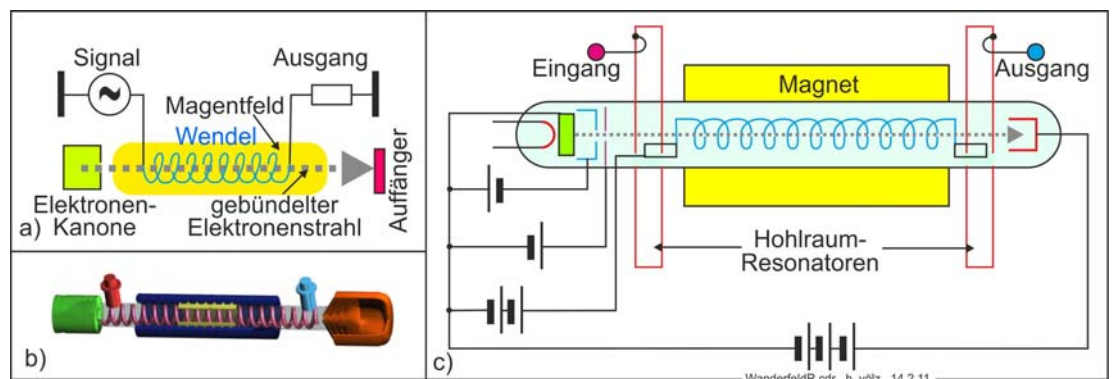
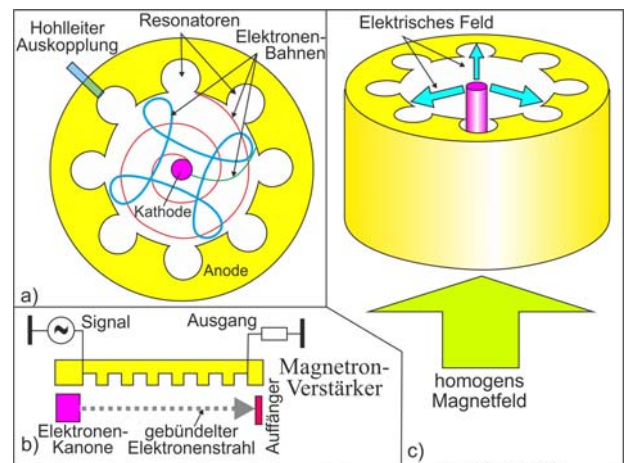


Bild 53. Wirkprinzip der Wanderfeldröhre (a), Struktur (b) sowie Aufbau und Schaltung (c).

Ein **Magnetron** nutzt ebenfalls Laufzeiteffekte, kann aber nur als Oszillator betrieben werden. Bei ihm sind die Hohlraum-Resonatoren als Schlitze oder Spalte, endlos kreisförmig in einem Zylinder angeordnet (**Bild 54**). Sie sind zugleich die geerdete Anode. Zentral befindet sich die Kathode, die mit einer negativen Spannung von einigen kV betrieben wird. In Richtung der Zylinderachse (c) wirkt ein homogenes Magnetfeld B (häufig ein Permanentmagnet). Deshalb ist auch die Bezeichnung Kreuzfeldröhre (crossfield amplifier, CFA) gebräuchlich. Das Zusammenwirken des radialen elektrischen und axialen magnetischen Feldes zwingt Elektronen auf Zykloidenbahnen. Nur Elektronen mit der „richtigen“ Geschwindigkeit erreichen die Resonatoren der Anode streifend (blau). Insgesamt bildet sich die Oszillatorfrequenz aus, bei der annähernd gilt $f_{Res} \approx 20 \text{ GHz} \cdot B$ (in Tesla). Typische Frequenzen sind 0,3 bis 30 GHz, der **Wirkungsgrad** ist mit bis zu 80 % sehr hoch. Dagegen ist die **Lebensdauer** mit 5000 h recht gering. Im Dauerbetrieb werden bis zu 30 kW, im Impulsbetrieb (u. a. Radar-Anlagen) bis 10 MW erreicht.

Bild 54. Zum Aufbau und Betrieb des Magnetrons.



²⁴ Griechisch *klýzein* anbranden; -tron Suffix zur Bezeichnung eines Gerätes, Werkzeugs.

Die häufigsten Anwendungen des Magnetrons erfolgen mit der Industriefrequenz von 2,45 GHz industriell (u. a. Trocknung ≤ 6 kW), im Haushalt (Mikrowellenherde ≈ 1 kW) und der Medizin (Erwärmung ≈ 200 W). Das Magnetron-Prinzip ist auf linear angeordnete Resonatoren (Runzelleiter) übertragbar (b). Dann entstehen auch Wanderfeld-Verstärker, u. a. das Amplitron.

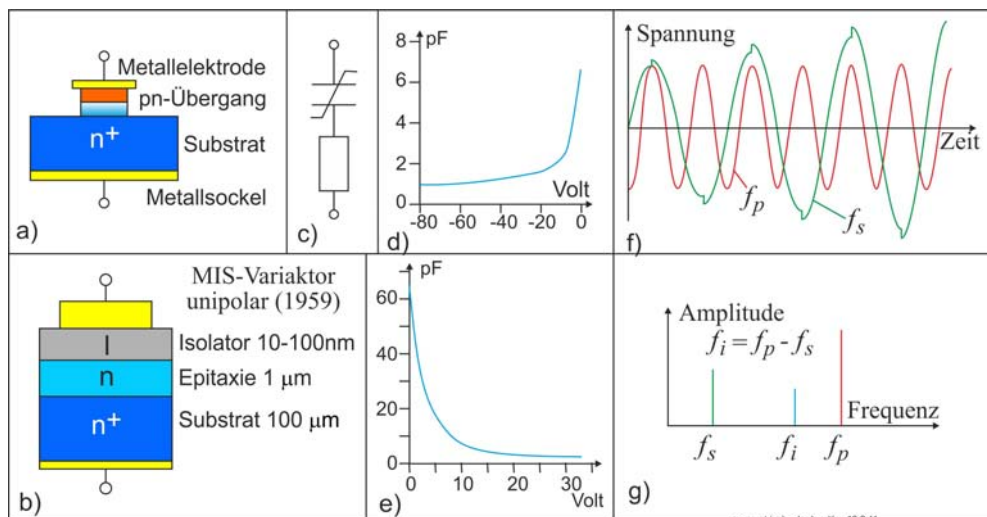
Beim **parametrischen Verstärker** werden besondere Bauelemente mit der doppelten Frequenz der eigentlichen Schwingung so geändert, dass zusätzliche Verstärkung auftritt. Das entspricht der Analogie auf einer Schaukel. An den Stellen des höchsten Ausschlags wird der Körperschwerpunkt durch in die Hockegehen bzw. sich Ausstrecken so verlagert, dass der Ausschlag zunimmt (**Bild 55**). Mit diesen Bewegungen allein kann jedoch die Schaukel nicht von der Ruhe in Bewegung gebracht werden.

Bild 55. Schaukel als Analogie für parametrische Verstärker.

In der Technik werden z. B. Reaktanz-Verstärker erzeugt. Dazu dienen vor allem pn- oder Impatt-Dioden, MIS-Variaktoren (Metall Isolation Silizium) oder Schottky-Kontakte. Dann wird eine Kapazität benutzt, deren Wert von der angelegten Spannung abhängt als Variaktor eingefügt. Bei-

spielen und Daten zeigt **Bild 56** a bis e. Den entsprechenden Zeitverlauf zeigt (f). Darin bedeuten f_s die zu verstärkende Signalfrequenz und deren doppelte Frequenz f_p zum „Pumpen“ als Idler²⁵- oder Hilfsfrequenz. Infolge der so wirksam werdenden Parameter-Änderungen wird die Signalfrequenz deutlich verstärkt. Da die Schaltung wie ein negativer Widerstand wirkt, steigt die Amplitude exponentiell mit der Zeit (g); mit Ähnlichkeit zur Pendelrückkopplung. Durch die nicht mehr exakten sinusförmigen Signalschwingung treten auch Vielfache der Summen- und Differenzfrequenzen auf: $f_x = n \cdot f_p \pm m \cdot f_s$. Technisch ist es besonders schwierig die Idler-Frequenz $f_i = f_p - f_s$ ausfiltern.

Bild 55. Beispiele und Schwingungen für parametrische Verstärker



Zur Beschreibung des Sekundärelektronenvervielfachers (SEV) ist von einem Einzelatom auszugehen. Es kann ein einzelnes, freies Atom nur dann verlassen, wenn ihm die Ionisierungsenergie W zugeführt wird. Mit der Elementarladung $e_0 \approx 1,602 \cdot 10^{-19}$ kann sie in Elektronenvolt (eV) umgerechnet werden $U = W/e_0$. Durch die Wechselwirkung der Atome im Festkörper ist aber die Austrittsarbeit deutlich geringer. Insgesamt können folgende Emissionsarten wirksam werden:

1. Feldemission, sie benötigt Feldstärken um 10^7 V/cm, was sehr kleine Radien also Spitzen verlangt, z. B. Feldelektronenmikroskop.
2. Stoßemission als Ionenstoß und Sekundäremission.
3. Fotoemission ausgelöst durch Lichtquanten
4. Thermoemission beruht auf der statistischen Wärmebewegung.

Bei der Stoßionisation besteht eine Ähnlichkeit zu Wassertropfen, der beim Auftreffen auf eine Wasseroberfläche noch mehr Wasser hochschleudert. Das Verhältnis von auftreffender zu austretender Elektronenzahl ist der Ausbeutefaktor δ . Er hängt auch von der Energie der Primärelektrons ab. Beispiele zeigt **Bild 56** (c) in linearer und (d) in logarithmischer Darstellung. Grob können drei Stoffarten unterschieden werden:

- $\delta < 1$, Alkali- und Erdalkalimetalle sowie Kohlenstoff (Graphit).
- $1 < \delta < 2$, die meisten Metalle.
- $\delta > 2$, Metalloxide und -chloride (BeO ≈ 10 , Mg ≈ 8 , NaCl ≈ 6), Metalllegierungen GaP-Cs ≈ 5 , Cu-Mg ≈ 13 , Ni-Be ≈ 12 , Al-Cu ≈ 10 .

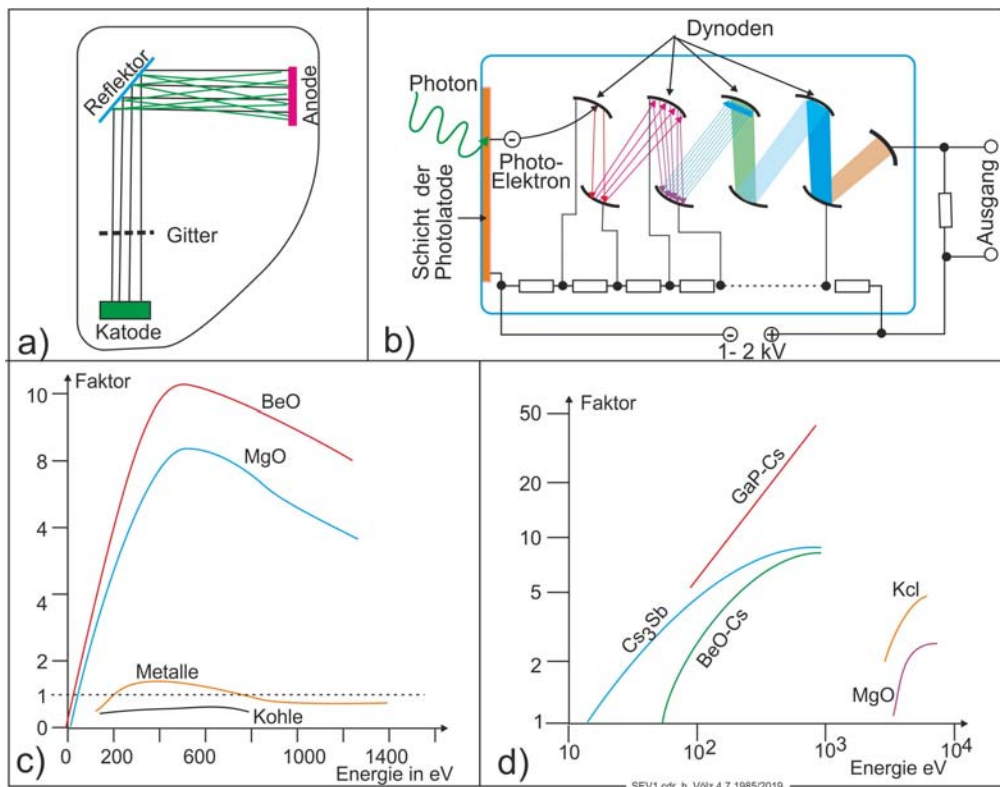


Bild 56. Varianten und Eigenschaften von Sekundärelektronenvervielfachern: (a) Elektronenröhre mit SEV, (b) Beispiel eines Fotovervielfachers, (c) Verstärkungsfaktor δ in linearer und (d) in logarithmischer Darstellung.

²⁵ englisch **idler** Müßiggänger

Die Verzögerung der Sekundärelektronen beträgt ca. 10^{-10} s, was die Verstärkung sehr hoher Signalfrequenzen erlaubt. Die SEV wurden nur Anfangs zusätzlich bei Elektronenröhren genutzt (a). Große Bedeutung erlangten sie als Fotovervielfacher (photomultiplier tube, PMT) mit einer Vakuumfotозelle als Katode (b) im evakuierten Glaskolben (10^{-6} bis 10^{-5} Pa). Die Reflektoren (Dynoden) werden mit Material von hohem δ beschichtet und an hohes, aber abgestuftes positives Potential angeschlossen. Der Verstärkungsfaktor wächst exponentiell mit der Anzahl der Dynoden. Typische Multiplier haben etwa 10 Dynoden. Werden an jeder Dynode 4 Elektronen pro auftreffendes Elektron herausgeschlagen, so entsteht eine rauscharme Verstärkung von $4^{10} \approx 10^6$ mit nur etwa 30 ns Verzögerung, was höchste Signalfrequenzen möglich macht. Wegen ihrer hohen Empfindlichkeit müssen die meisten Fotovervielfacher bei Betrieb vor Tageslicht geschützt werden. Der Lichteinfall könnte zu viele Photonen und damit einen zu hohen Strom erzeugen, was die Beschichtung der Dynoden irreversibel schwächen kann („Erblinden“) oder sogar zu einem Durchbrennen führen kann. Es gibt mehrere Variationen des SEV-Prinzips. So wird es ähnlich bei Szintillations- und Tscherenkow-Zählern für sehr geringe radioaktive Strahlung verwendet.

6. O-Verstärker für Bilder und Schall

Der sehr umfangreiche Inhalt dieses Gebietes ist nicht unmittelbar einsichtig und wurde bisher auch kaum unter dem Gesichtspunkt der Verstärkung behandelt. Dennoch sprechen viele Fakten für diese Zusammenfassung und vermehren so die Universalität von Verstärkung. Daher sei hier zunächst – bevor ausgewählte konkrete Beispiele behandelt werden – ein grober Überblick gegeben.

Primär betreffen die Techniken der O-Verstärker die zweidimensionale Bildwahrnehmung und -darstellung. Dabei hierbei aber hauptsächlich die Größe zwischen bestimmten Punkten des Bildes betrachtet wird, ist der Begriff Vergrößerung üblich. Dennoch betrifft der Inhalt ziemlich genau den der Verstärkung bei anderen Ausprägungen, wie Signalgröße, Kraft usw. In der Optik leisten also Vergrößerungen das Äquivalente. Das sind zunächst und vor allem Linsen, Brillen, Hohlspiegel, Fernrohre und Mikroskope. Teilweise bewirken sie das auch im Infrarot, Ultraviolett und bei Röntgenstrahlung. Bei den Mikroskopen gehören sogar die elektronischen Varianten, einschließlich der rasterelektronischen Abtastungen dazu. Sonderformen betreffen ähnlich auch elektronische Displays, Projektoren, Beamer Headset sowie elektronische Bildvergrößerungen insbesondere für Sehbehinderte. Teilweise erfolgt bei ihnen ein zeitliches Abtasten (Scannens). Zusätzliche und nachträgliche Bildverbesserungen (Kontrast, Farbe, Schärfe usw.) erfolgen mit leistungsfähigen Bildbearbeitungsprogrammen. Sonderfälle bei der Optik sind die Möglichkeiten von Raumwahrnehmung, also Stereoverfahren und Holografie. Für die Astronomie (Weltraum) sind auch elektromagnetische Wellen, insbesondere mit Spiegeln und sogar umfangreiche Antennenfelder wichtig. Doch hierbei handelt es sich eigentlich ähnlich wie der Fotografie von Landschaften, Gebäuden usw. eigentlich um Verkleinerungen auf eine unserem Auge angepasste Verkleinerung. Dennoch ist es aber indirekt eine Verstärkung für unser Sehen. Im erweiterten Sinne gehören dann auch alle Radar und Verwandtes zu den O-Verstärkern und folglich auch die vielfältigen Antennen der Rundfunk- und Fernsehertechnik. Speziell sind dabei vor allem die verschiedenen Methoden mit Richtwirkung mit der dazu gehörenden Reziprozität von Senden und Empfangen entscheidend. Mittels der Richtwirkung wird auch die Energiebündelung eine energetische Verstärkung wirksam, das beginnt im Optischen stark vereinfacht beim Brennpunkt.

Viele der genannten Aspekte sind analog auf Schall zu übertragen. Primär betrifft es hier die Vielfalt der Mikrophone und Lautsprecher, einschließlich der Richtwirkung verschiedener Zusammenschaltungen (wie Lautsprecherfelder), aber auch von Zusätze wie Schallwände, Bassreflex usw., sowie die besonderen Anwendungen, z. B. bei Megaphonen, Hörgeräten sowie Resonanzböden bei Musikinstrumenten. Beim Ultra- und Infraschall (Kommunikation der Fledermäuse und Wale) sind auch Frequenzumsetzungen für das Hörbar-machen nützlich. Sondergebiete sind noch Raumakustik und Ultraschallreinigung.

Da elektromagnetische Wellen (vereinfacht als Licht bezeichnet) und Schall mehrfach in die weiteren Betrachtungen eingehen, zeigt **Bild 57** einen Überblick und Vergleich der beiden Frequenzbereiche. Infolge der sehr unterschiedlichen Ausbreitungsgeschwindigkeiten – Licht $\approx 3 \cdot 10^8$ m/s, Schall in Luft ≈ 440 m/s – musste die Skala für die Wellenlängen des Schalls ergänzend links oben eingezeichnet werden. Im unteren Bildteil sind die Anwendungsbereiche für verschiedene Technologien zur Beeinflussung von Licht ergänzt.

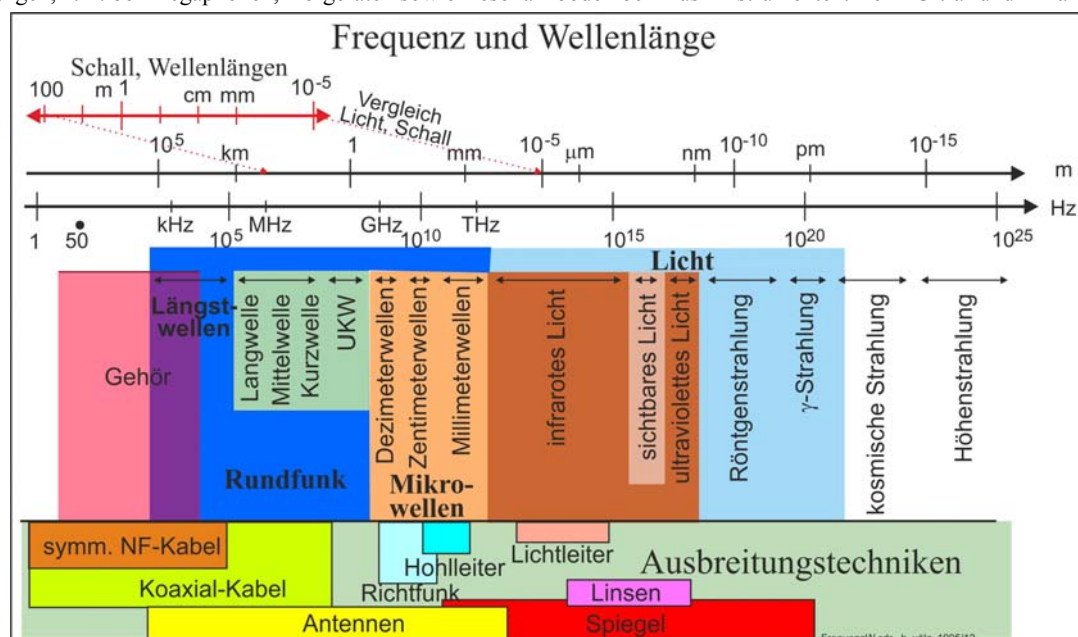
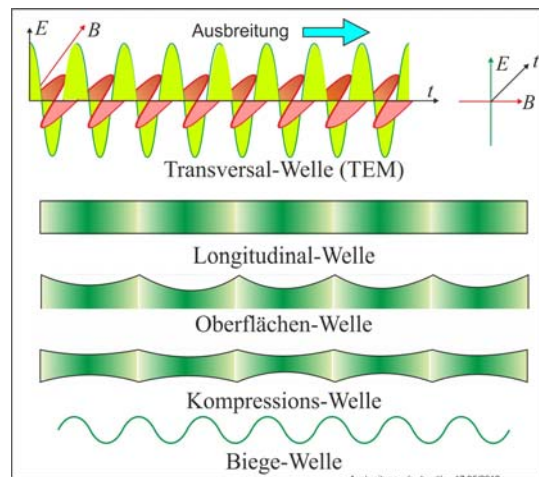


Bild 57. Frequenzen und Wellenlängen bei elektromagnetischen Wellen und Schall.

Es gibt unterschiedliche Eigenschaften der Wellen. Alle existierenden Varianten sind in **Bild 58** gezeigt. Bei den elektromagnetischen Schwingungen stehen die elektrischen und magnetischen Felder immer senkrecht zueinander und senkrecht zur Ausbreitungsrichtung. Ihre Amplitudemaxima wechseln sich nacheinander ab (Transversalwelle). Schall breitet sich in der Luft dagegen als Longitudinal-Welle aus. Dabei entstehen nacheinander Verdichtungen und Verdünnungen der Luftmoleküle. Infolge dieser Unterschiede sind bei Licht und Schall zumindest teilweise andere Verfahren zur Beeinflussung erforderlich. Bei einer Flüssigkeit können auch Oberflächen-Wellen auftreten. In anderen Materialien auch noch Kompressionswellen mit Verdichtungs- und doppelseitige Oberflächenänderungen. Bei Seilen, Drähten usw. gibt es schließlich noch die Biege-Wellen.

Bild 58. Typischen Eigenschaften aller möglichen Wellenausbreitungen.



Vielfalt optischer Linsen

Die Beschreibung des optischen Geschehens kann theoretisch gleichwertig mit der Methode des Lichtstrahls oder der Wellenfront erfolgen (**Bild 59**). Bei relativ einfachen Vorgängen sind meist die Lichtstrahlen vorteilhafter (a), jedoch bei der Holographie, Interferenzen usw. sind es die Wellenfronten (b). Das ist jedoch meist schwierig zu verstehen (s. u.)

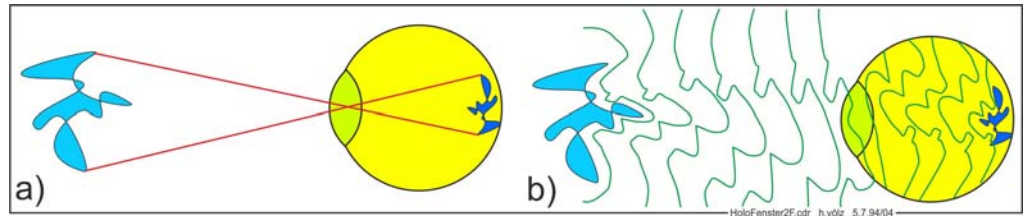
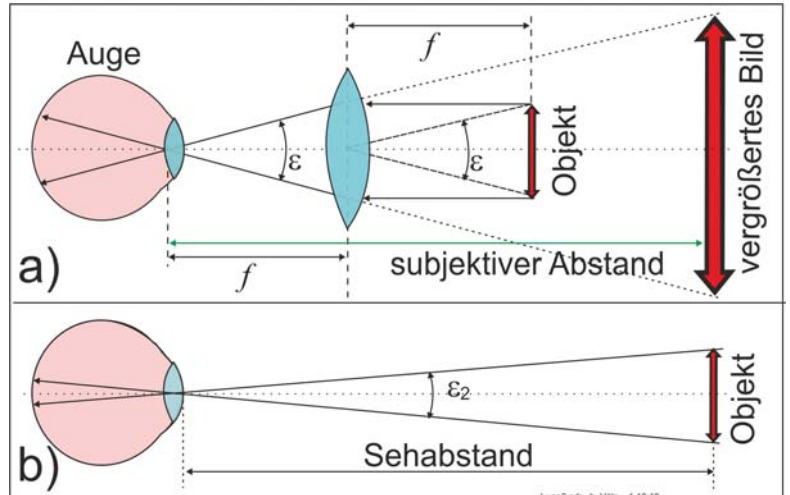


Bild 59. Die beiden Hauptverfahren für optisches Geschehen.

Linsen sind meist kreisförmig gestaltete und optisch durchlässige Gebilde mit vorwiegend kugelförmig gekrümmten Oberflächen. Denn nur sie lassen sich mit einfachen Technologien exakt herstellen (schleifen). Eine sehr gebräuchliche Anwendung für eine vergrößerte Betrachtung von etwas stellt die **Lupe** dar. **Bild 60** veranschaulicht die dazu gehörende Wirkungsweise (a). Die Linse wird dazu ungefähr im Abstand der Brennweite f vom Auge gehalten. Vom Objekt aus, das sich ebenfalls etwa im Abstand f befindet, gehen dann die Lichtstrahlen parallel zur Linse hin und werden dort zueinander umgelenkt. So entstehen die Öffnungswinkel ε . Werden nun die Lichtstrahlen zwischen Auge und Linse darüber hinaus weiter verfolgt, so entsteht im optimalen Abstand (auf seine Berechnung sei hier verzichtet) das virtuelle, aber subjektiv sichtbare Bild des vergrößerten Objektes. Würde sich das Objekt ohne Linse in diesem Abstand befinden, so entstünde ε_2 und das subjektiv kleinere Objekt.

Bild 60. Strahlengang (a) und Auswirkungen (b) bei einer Lupe.



Die **geschliffene Glaslinse** ist nur eine Variante aller möglichen Linsen. Ihre Vielfalt demonstriert **Bild 61**. Bei der Lupe in Bild 60 wurde eine typische **Sammellinse** benutzt. Gemäß Bild 61a ist sie meist bikonvex, also mit beidseitiger Mittendickung. Parallele Lichtstrahlen die zentral auf sie treffen, werden in einem Brennpunkt f gebündelt und divergieren danach wieder. Eine **Zerstreuungslinse** (b) ist dagegen bikonkav, also im Mittelteil dünner. Auf sie treffende parallele Lichtstrahlen divergieren danach. Jedoch in deren Rückverfolgung existiert ein virtueller Brennpunkt f' , von dem das Licht scheinbar herzukommen scheint. Eine technologische Abwandlung der Sammellinse ist die **Fresnel-Linse**. Formal kann sie aus einer nur einseitig konvexen Sammellinse entstanden gedacht werden (c). Von ihr wird der rote Teil kreisförmig herausgefräst. Dann wird der blaue Rest flach zusammengedrückt. Technologisch wird jedoch die recht unregelmäßig erscheinende Oberfläche in eine gut durchsichtige Plastikfolie gepresst. So entsteht ein flaches Blättchen, das genauso wie eine Sammellinse wirkt und daher gut mit Scheckkarten usw. aufbewahrt werden kann. Es gibt solche sehr handlichen Linsen bis etwa A5-Größe. Notwendig ist dabei nur, das weitaus mehr Stufen als im Bild 61c erzeugt werden.

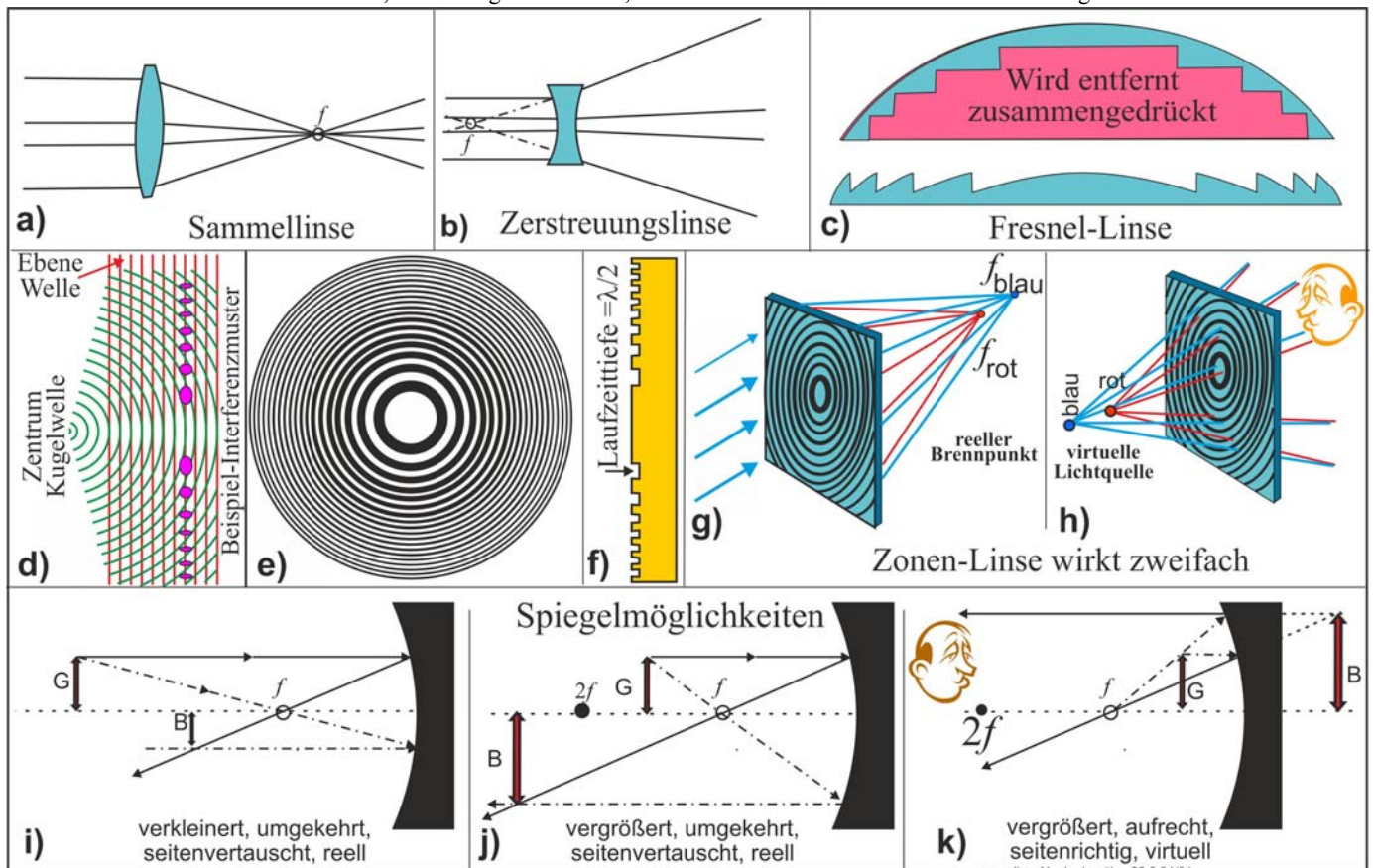


Bild 61. Die vielfältigen Linsen und der Hohlspiegel.

Völlig anderer Art ist aus der Holografie stammende **Zonenlinse**, (manchmal leider mit der Fresnellinse verwechselt). Ihre Beschreibung ist am besten mit den Wellenfronten von Bild 59b und dann mittels Bild 61 (d bis f) zu erklären. Nach (d) bereitet sich eine Kugelwelle von einem Punkt allseitig aus (blau). Sie interferiert mit einer streng ebenen Welle (rot), die scheinbar aus dem Unendlichen kommt. In der Überlagerung entstehen dann Punkte (Linien), wo sich ihr Licht addiert und andere wo es sich aufhebt. Ausgewählte Licht addierende Punkte sind im Bild lila hervorgehoben. Werden sie der zum Bild senkrechten Ebene zusammengefasst, so ergibt sich das Aussehen der Zonenlinse (e). Für ihre Anwendung müssen die schwarzen Kreise undurchsichtig sein. Das lässt sich gut mit flachen Metallringen realisieren, die dann aber mit wenigen Stegen stabil in ihrer Lage gehalten werden. Auch hier ist eine Pressung in Kunststoff möglich. Gemäß (f) müssen dann die Ringe um $\lambda/2$ vertieft eingepresst werden. Dann tritt dort eine Interferenz auf, die den Lichtdurchtritt verhindert.

Die Wirkung der Zonenlinse weicht aber deutlich von den zuvor besprochenen Linsen ab. Je nach Anwendung wirkt sie als Sammellinse (g) bei parallelem Lichteinfall oder als Zerstreuungslinse (h) bei rückwärtiger Beobachtung, vgl. (a) und (h). Außerdem besitzt sie für jede Wellenlänge einen anderen (reellen bzw. virtuellen) Brennpunkt. Das hat aber auch Vorteile: Bei einer metallischen Ausführung der Ringe ist sie in allen Frequenzbereichen von Längstwellen bis zu Röntgenstrahlen anwendbar. So wurde sie z. B. beim ROSAT-Satelliten zur Erfassung der Röntgenstrahlung im All benutzt. Auch Schallanwendungen sind möglich.

Ähnlich wie die einfachen Linsen (a und b) können auch (parabelförmig) gekrümmte **Spiegel**-Flächen wirken (i bis k). Im Bild wird das nur für praktisch wichtigen Hohlspiegel angenommen. Er entspricht näherungsweise einer Sammellinse. Jedoch gegenüber der Linse wird das Licht reflektiert und die Abbildungen hinter der Linse werden rückwirkend in Richtung des ankommenden Lichtstrahls erzeugt. So entstehen die drei ausgewählten und besonders wichtigen Varianten. Je nach Lage des Gegenstands **G** in Bezug auf den Brennpunkt entstehen die drei typischen Abbildungsvarianten für seine Abbildung **B**, deren Eigenschaften darunter hinzugefügt sind. Besonders nützlich ist (k) mit einer virtuellen vergrößerten Abbildung, so wird u. a. der Stirnspiegel bei der Untersuchung des Kehlkopfs benutzt. Weitere Varianten werden noch beim Spiegelteleskop (s. u.) behandelt.

Eine besonders vielfältige Anwendung von Linsen erfolgt bei der Fotografie. Ihr historischer Ursprung war die Lochkamera, deren Wirkung **Bild 62c** schematisch zeigt. Im total verdunkelten Raum befindet sich nur ein kleines Loch, durch das sich die draußen hell beleuchtete Umwelt an der gegenüberliegenden Wand spiegelbildlich seitlich verkehrt und kopfstehend abbildet. Genau dieses Abbildungsprinzip leistet – aber deutlich effektiver – eine Linse (a). Mit der Brennweite f und der Gegenstandsweite g des Objektes **O** entsteht das Abbild **B** in der Bildweite b . Für deren Größenverhältnis gilt zusätzlich zur angegebenen Formel $O/B = g/b$.

Für die weiteren Betrachtungen ist der Durchmesser D der Linse wichtig (b). Er wird auf mehrfache Weise benannt. Mit dem Brechungsindex n folgt die **Numerische Apertur**²⁶ $N_A = n \cdot \sin(\varphi)$. Die **relative Öffnung** ist das Verhältnis $\tilde{O} = D/f$. In der Fotografie (s. u.) wird sie – z. B. mittels einer **Irisblende** – eingeschränkt ($\tilde{O}_B \leq \tilde{O}$ bzw. $D_B \leq D$). Die sich dann ergebende Öffnung heißt **Blende B** und wird als Bruch angegeben (z. B. 1:2 bis 1:32). Reziprok ist der **Blendenwert** $k = 1/B = D_B/f > 1$. Teilweise gelten diese Werte exakt nur dünne Linsen mit der Dicke $d \ll r$ und $r = D/2$.

Leider entstehen bei einer optischen Abbildung auch Mängel. Einige zeigt **Bild 63**. Gemäß der Strahlenoptik bestimmt die Größe des Loches bzw. der Blende die Schärfe eines Bildes. Mit kleinerer Blende werden die Abbildungen zwar immer dunkler, dafür aber immer schärfer. Doch die Zunahme der **Schärfe** wird schließlich durch die Wellenoptik begrenzt. Bei einem sehr kleinen Loch tritt zusätzlich eine kreisförmige Beugungsfigur auf. Folglich existiert das schärfste Bild bei einer optimalen Blende bzw. Lochgröße. in **Bild 63a** verdeutlicht das die Kurve und die Abbildung einer leuchtenden Glühlampenwendel.

Mit der Blende ändert sich auch der Entfernungsbereich des Objektes der scharf abgebildet werden kann. Er wird die **Schärfentiefe** genannt, früher auch Tiefenschärfe dafür gebräuchlich. Mit der elektronischen Fotografie ist durch die nun verwendeten kurzen Brennweiten der Objektive dieser Einfluss fast nicht mehr feststellbar. Doch voll erhalten geblieben sind die anderen Mängel. Mit (b) sei auf die **Farbsäume** hingewiesen, die vor allen an kontrastreichen Grenzen auftreten Sie werden dadurch hervorgerufen, dass einfache Linsen ähnlich – wenn auch wesentlich geringere **Farbfehler** ähnlich wie bei den Zonenlinsen in Bild 61 (g und h) auftreten, Sie lassen sich durch Optiken mit komplexen Linsenaufbau (s. u.) weitgehend reduzieren. Auch **Geometriefehler** treten auf. Bild 63c zeigt Beispiele für perspektivische und kissenförmige Verzerrungen. Für weitere Fehler sei auf Spezialliteratur verwiesen u. a. [Völ05].

Die verschiedenen Abbildungsmängel (Fehler) einer einfachen Linse sind weitgehend durch eine Kombination aus mehreren Linsen mit verschiedenen Glassorten und Geometrien erheblich zu reduzieren. Dazu besitzen die einzelnen Glassorten abweichende Verläufe der Lichtbrechung bei den einzelnen Spektralfarben. Geometriefehler reduziert vor allem die Kombination von Sammell- und Zerstreuungslinsen. Beispiele einiger bewährter Linsenaufbauten zeigt **Bild 64**. So werden auch Optiken erzeugt, die als Teleobjektive eine besonders lange Brennweite für die Aufnahme weit entfernter Objekte besitzen. Die entgegengesetzte Variante sind Fischaugenoptiken (Weitwinkel f und g), die einen großen Winkelbereich abbilden.

Fotografie

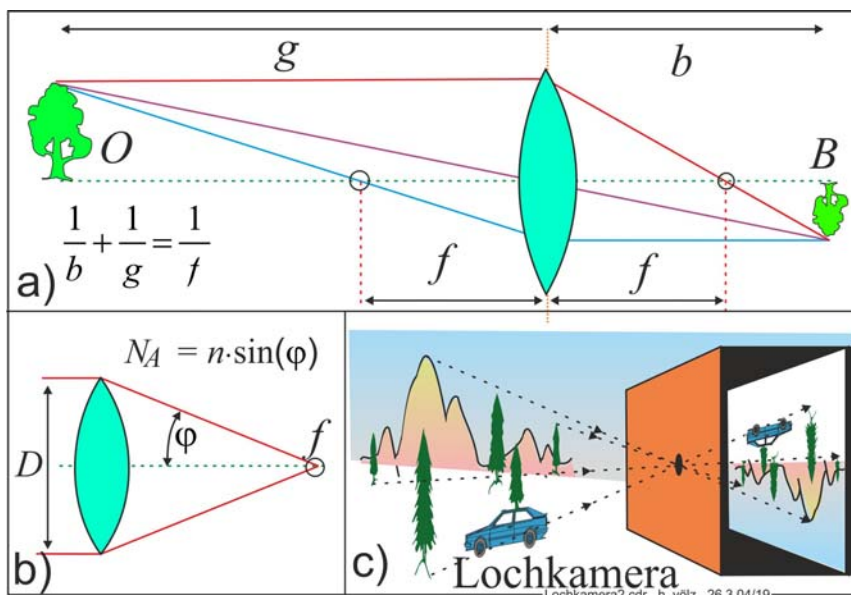


Bild 62. Von Lochkamera bis Fotografie.

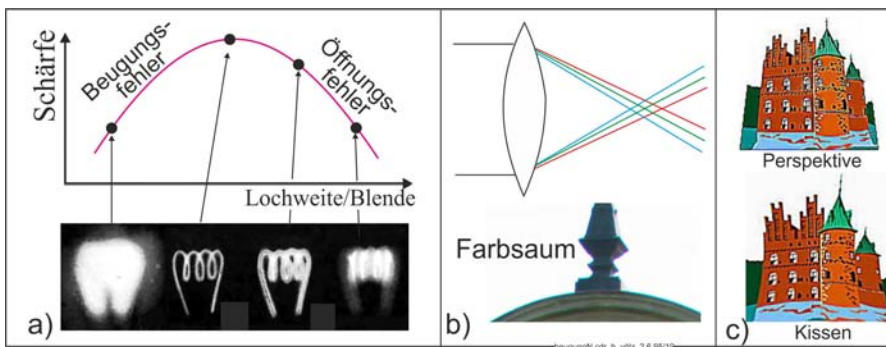


Bild 63. Darstellung besonders wichtiger Abbildungsfehler.

²⁶ Apertur von lateinisch aperire öffnen, apertura Öffnung

Mit sehr umfangreichen Berechnungen entstehen seit geraumer Zeit auch Zoomobjektive (Variooptiken) mit frei wählbarer Brennweite. Im Beispiel (h) wird dazu lediglich ein Teil des Objektivs vor- bzw. zurückgeschoben (h). Schließlich zeigt (i) noch ein Spiegelobjektiv aus zwei Hohlspiegeln und mehreren Linsen. Es besitzt eine sehr große Brennweite bei beachtenswert großer Öffnung.

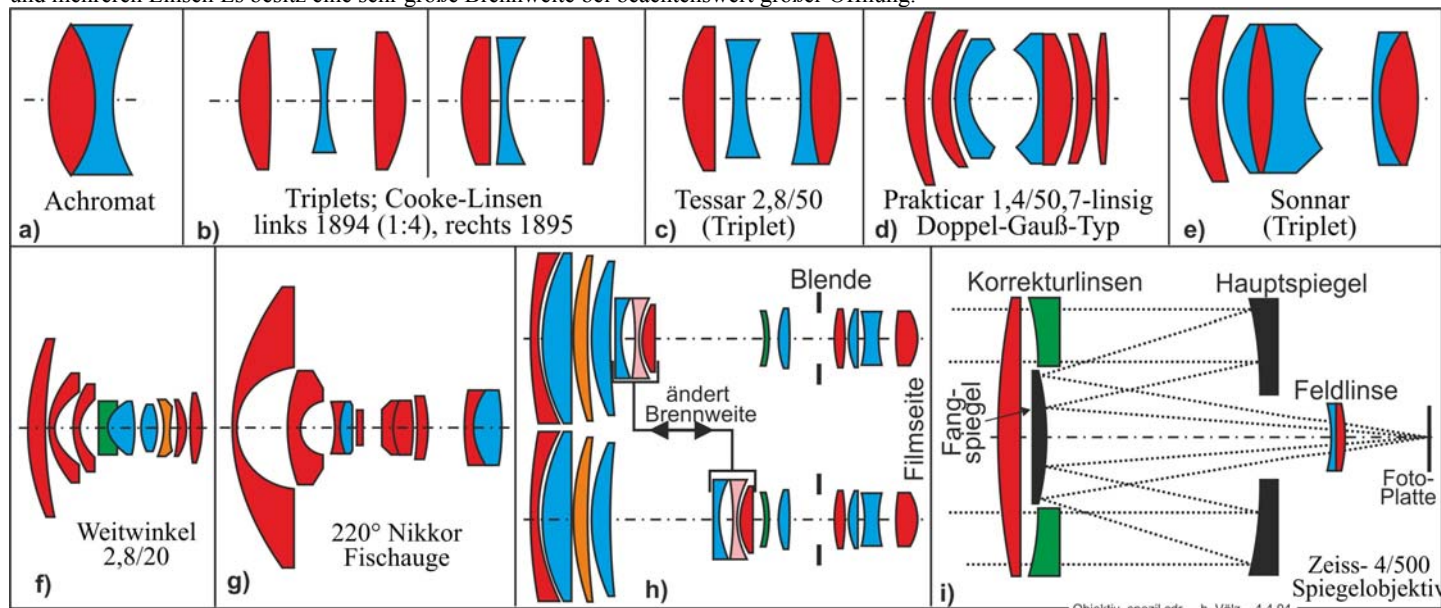
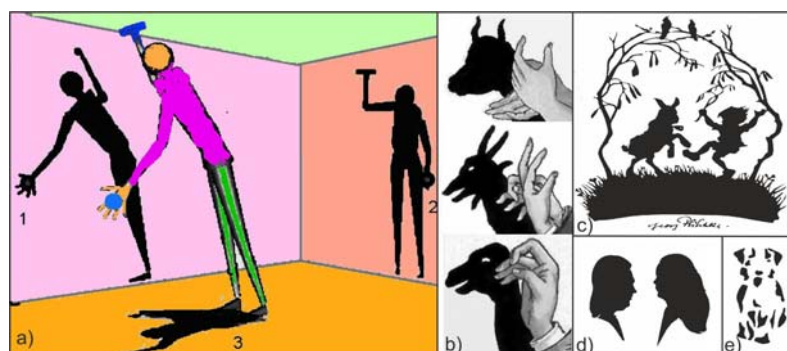


Bild 64. Beispiele für den komplexen Aufbau leistungsfähiger Fotoobjektive.



Mit der Fotografie entfernt verwandt – wenn auch stark vereinfachend – ist das **Schattenbild**. Es hat aber auch Vorteile. So können wie im **Bild 65a** durch unterschiedliche Lichtquellen leicht mehrere Perspektiven parallel gewonnen werden. Es lassen sich Schattenspiele veranstalten (b). Dann gibt es die bekannten Scherenbilder nach Plischka (c) und sogar Portraits (d). Schließlich sein noch jene Bilder genannt, die beim Sprühen auf Straßen und Häusern erzeugt werden (e).

Bild 65. Beispiel für mit drei Lichtquellen erzeugte Schattenbilder (a) und weitere Schattenanwendungen.

Fernrohre und Teleskope

Mit der Lupe sind Vergrößerungen von 2-fach bis maximal 10-fach aber nur für nahe Objekte möglich. Die Fotoobjektive bilden auch weiter Entfernertes ab, jedoch nicht allzu groß. Deutlich mehr leisten jedoch unterschiedliche Fernrohre. Die wichtigsten Varianten zeigt **Bild 66**. Besonders einfach ist das **Galilei-Fernrohr** (a). Es benötigt nur eine Sammellinse als Objektiv und eine Zerstreuungslinse als Okular. Die Sammellinse bereitet nur die Abbildung des weit entfernten Gegenstandes seiten- und höhenvertauscht vor. Ihr Brennpunkt liegt dabei etwa im Auge. Die Zerstreuungslinse erzeugt dann das scharfe und richtig stehende Bild. Das einfache Prinzip ist preiswert herzustellen, ist jedoch nicht sehr leistungsfähig. Es wird daher bevorzugt für wenig anspruchsvolle Anwendungen eingesetzt. Das sind vor allem Opern-, Theater- oder Museumsgläser, die auch teilweise einfach Fernglas und bei höheren Leistungen **Feldstecher** genannt werden. Die üblichen Vergrößerungen betragen bei Operngläsern 3-fach, für Feldstecher 6- bis 10-fach. Das Prinzip des Galilei-Fernrohrs hat den Vorteil, dass es auch mit geringem Gewicht (bis etwa 300 g gegenüber sonst bis zu 1,2 kg) und teilweise gut zusammenlegbar für einen Transport als Kompakt- oder Taschenfernglas ausführbar ist. Bei einäugiger Bauweise heißt das Galilei-Fernrohr auch **Monokular**. Eine Spezialanwendung ist das **Zielfernrohr**.

Das **Kepler-Fernrohr** benutzt im einfachsten Fall nur zwei Sammellinsen (b). Das Objektiv erzeugt dabei ein reelles Zwischenbild, das aber doppelt verkehrt, d. h. auf dem Kopf stehend und seitenverkehrt, also um 180° gedreht ist. Mit dem Okular in der Funktion einer Lupe wird es vergrößert betrachtet. Zur Anwendung wird das Zwischenbild mit einem Umkehrprisma (c) zurückgedreht. Zur Vereinfachung der Darstellung ist es nicht in (b) eingezeichnet. Die Leistungsfähigkeit des Kepler-Fernrohres ist dem Galilei-Fernrohr deutlich überlegen. Dafür werden dann insbesondere bei Großferngläsern beim Objektiv zwei bis fünf Linsen, beim Umkehrprisma zwei oder drei einfache Prismen und beim Okular meist drei bis sechs Linsen eingesetzt. Die Objektivöffnungen gehen oft deutlich über 50 mm hinaus. Insgesamt erlaubt das Prisma eine kompakte, wenn auch schwerere Gestaltung. Es existieren sowohl gut handhabbare als auch Großferngläser die mit Stativen benutzt werden. Ferner gibt es Zoomferngläser so genannte **Spektive** mit ausziehbaren Tubus, bei denen die Vergrößerung von 15- bis 50-fach frei einstellbar ist.

Sonderausführungen des Kepler-Fernrohres sind die **Periskope** (e). Im Beispiel erfolgt die Bildumkehr durch zusätzliche Linsen. Die bis zu mehreren Metern langen senkrechten Tubusse (Objektarme) lassen den Beobachter unsichtbar bleiben. Das ist z. B. in Schützengräben und beim getauchten U-Boot nützlich. Mit großen Tubus-Abständen ist große Tiefensicht möglich. Es gibt auch hier unterschiedliche Linsenanordnungen.

Ein Beispiel für den Übergang zu **Teleskopen**²⁷ mit Spiegeln (seit etwa 1900 als Glasguss) zeigt (d). Hiermit sind besonders große Eingangslinsen mit sehr langer Brennweite nutzbar. Sie werden jedoch vor allem in der Astrophysik als Großteleskope benutzt. Dabei existieren mehrere Varianten. Drei wichtige Ausführungen zeigt **Bild 66**. Als **Doppelrefraktoren** werden sie vielfach sowohl zur Beobachtung als auch zur Fotografie benutzt. Sie wurden ab 1893/99 mit 80/60 cm Apertur in Paris und Potsdam benutzt. In den USA folgten dann an der Lick- und Yerkes-Sternwarte Teleskope mit Objektiven von 91 und 102 cm. Damit war die optisch-mechanische Grenze des Linsen-Durchmessers und damit der Lichtstärke – infolge ihrer störenden Durchbiegungen – erreicht und es folgten die **Spiegelteleskope**, beginnend 1947 beim 5-Meter-Teleskop von Mount Palomar. Der größte war dann der 1975 gegossene 6-Meter-Spiegel für das russische Selentschuk. In den 1980er-Jahren folgten noch größere aber **segmentierte Spiegel**. Für etwa 2020 plant die ESO das 30-Meter-„European Extremely Large Telescope“ (E-ELT). Eine Weiterführung sind Felder mit

²⁷ Der Begriff Telskop wurde erst in der Neuzeit gebildet, von griechisch téle fern und skopéin beobachten, ausspähen. Es gab ihn aber schon im Altgriechischen als teleskópous weithin schauend

vielen elektronisch verkoppelten Teleskopen. Als Sonderausführung gibt es auch Teleskope mit rotierenden und daher parabelförmigen **Flüssigkeitsspiegel**. Seit geraumer Zeit existieren ferner Teleskope für **Infrarot, Ultraviolett, Radio- und Röntgenstrahlung**.

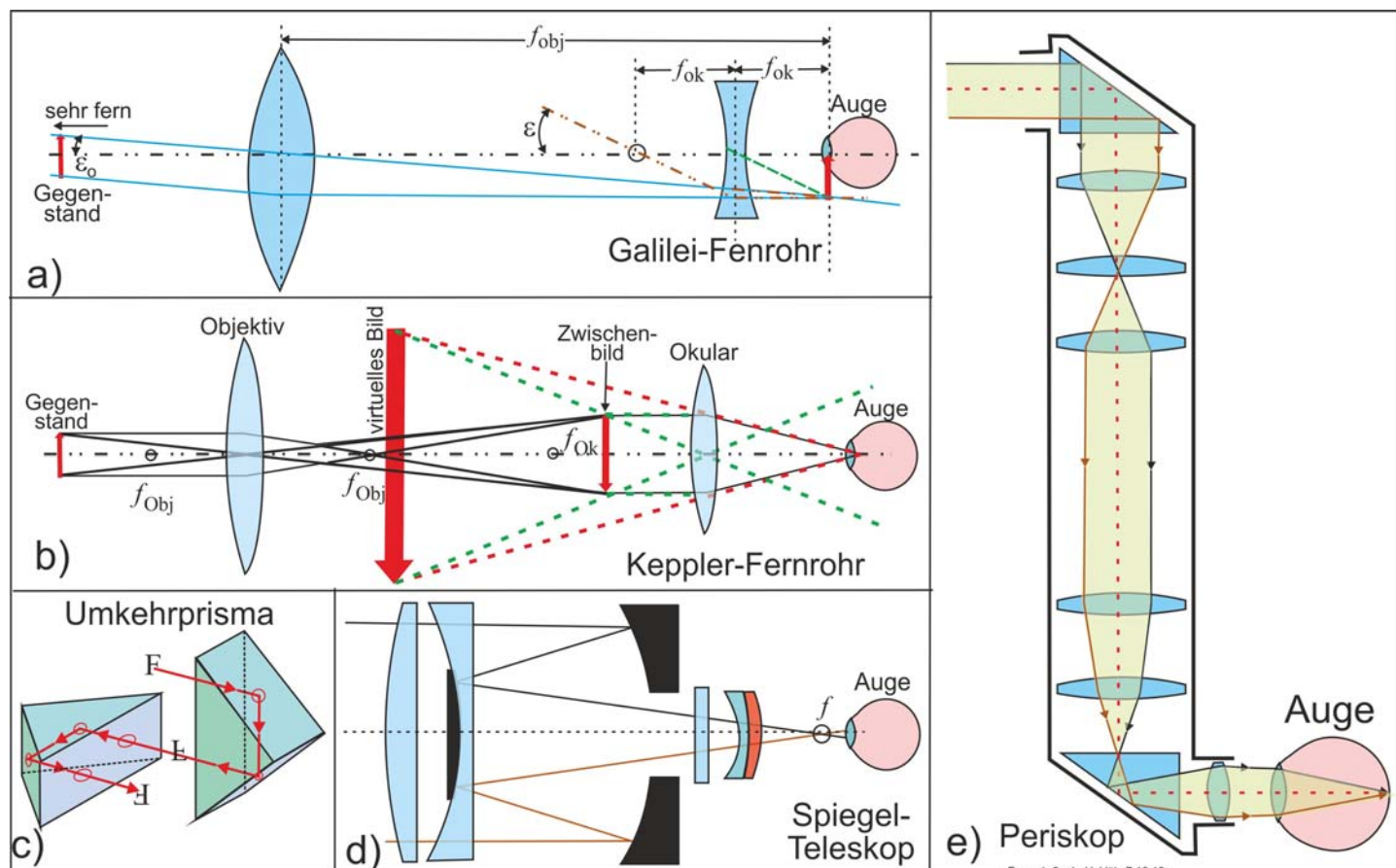


Bild 66. Versionen zum Oberbegriff Fernglas.

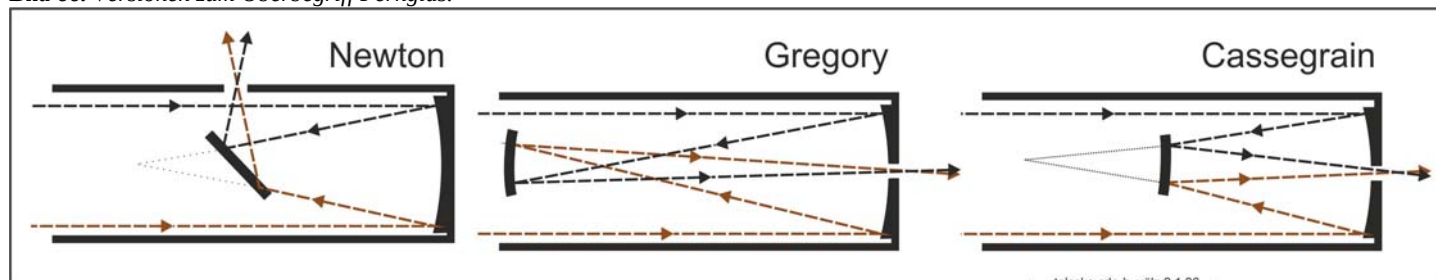


Bild 67. Einige einfache Varianten von Spiegelteleskopen.

Für das Weltall sind irdische und satellitengestützte Beobachtungen zu unterscheiden. Die Grenzen der irdischen Beobachtungen sind entsprechend der Erdatmosphäre auf zwei Frequenzbereiche eingeschränkt (Bild 68). Dabei ist das **Fenster** des optischen Sichtbereiches um Wellenlängen von etwa 300 (Begrenzung durch die Ozonschicht) bis 1 μm (Begrenzung durch Wasserdampf) begrenzt. Es gibt allerdings noch einige, sehr schmalbandige Fenster bis etwa 20 μm . Das Radiofenster reicht von wenigen cm bis zu etwa 100 m. Hierin liegt unter anderem auch die wichtige 21-cm-HI-Linie von neutralen Wasserstoffatomen, mit 21 cm bzw. 1420,4058 MHz

Da die Möglichkeiten des optischen Fenster bereits oben aufgezeigt sind, muss hier nur noch einiges zum Radiofensters ergänzt werden. Seit langem werden hier **Radio-Spiegelteleskope** (c) benutzt. Da auch hier konstruktive Grenzen bestehen, wurden schließlich in großen **passenden Tälern** feste Reflektorspiegel fest eingebaut (b). Durch die Verlagerung des Empfängers mittels Steuerung seiner Lage von den Seilen der hohen Masten, lassen sich unterschiedliche Sichtwinkel einstellen. So entstand auch in Arecibo (Hafenstadt 15 km von Puerto Rico) ein Teleskop ähnlich (b) mit 304 m Durchmesser). Seine Einweihung erfolgte 1963 und seit etwa 1992 werden die Messergebnisse vom SETI (search for Extraterrestrial Intelligence) der NASA durch viele freiwillige Teilnehmer im Internet ausgewertet.

Später entstanden noch erheblich leistungsfähigere **Antennenfelder** ähnlich (e). Durch Interferenz als VLBI (very long baseline interferometry, Langbasis-Interferometrie) mit Abständen bis zu 1000 km können ihre Signale selbst bei den großen Entfernungen elektronisch **per Computerrechnung zusammengefasst** werden (d). Ein Beispiel ist das USA.NM.VeryLargeArray. Schließlich wurden ähnlich auch die aus der UKW- und Fernsehtechnik bekannten **Yagi-Antennen** (s.u.) in großen Feldern ähnlich zusammengefasst, siehe Bild rechts, s. a. S. 37 ff.

Die **Satellitentechnik** kennt die Begrenzung der Wellenlängen nicht. Dabei ist allerdings zuweilen auch der Gebrauch von **Zonen-Linsen** erforderlich, vgl. Bild 61e. Der Einsatz von Satelliten erfordert aber einen beträchtlichen höheren finanziellen und technischen Aufwand. Außerdem ist die Steuerung aufwändig und die Ergebnisse sind nur vergleichsweise sehr langsam zur Erde zu übertragen.



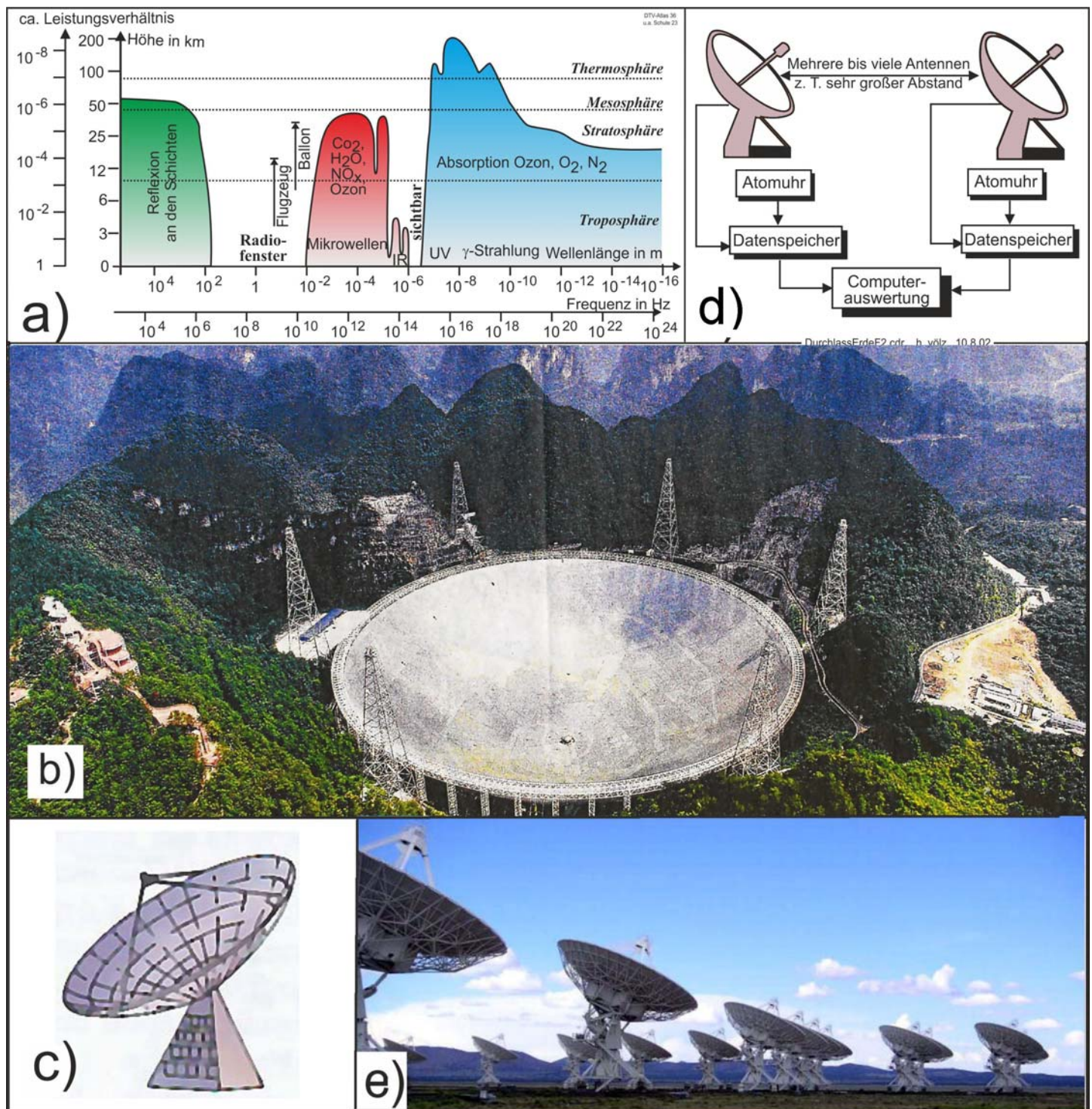


Bild 69. Beispiele für die Radioastronomie.

Mikroskope

Während Fernrohre und Teleskopie für die wesentlich erweiterte und mehr detaillierte Betrachtung entfernter Objekte ausgelegt sind, ermöglichen es Lichtmikroskope kleine Details von nahen Objekten zu analysieren. In diesem Sinne sind sie eine wesentliche Erweiterung der Möglichkeiten von Lupen. Die älteste Technik entstand um 1600 vermutlich in den Niederlanden. Anfang des 17. Jahrhunderts erhielt das mit Objektiv und Okular ausgestattete Mikroskop in Anlehnung an das Wort „Teleskop“ den Namen Mikroskop. Seinen Aufbau und die dazugehörigen Strahlengänge zeigt Bild 70a. Mit dem Objektiv wird vom Objekt ein deutlich vergrößertes und reelles Zwischenbild erzeugt. Es wird dann mit dem Okular (ähnlich wie bei der Lupe) nochmals vergrößert betrachtet. Durch das Wechseln des Objektivs und/oder Okulars sind Vergrößerungen von etwa 10- bis fast 1000-fach möglich. Die Grenze der Auflösung ist dabei infolge der Lichtwellenlänge auf etwa 0,2 μm begrenzt. Sie wird als Abbe-Limit bezeichnet, denn Ernst Abbe beschrieb die entsprechenden Gesetzmäßigkeiten Ende des 19. Jhs. Seit dem ist mit komplexen Interferenzmethoden und Beleuchtungstechniken noch eine gewisse Steigerung möglich.

Eine wesentlich höhere Auflösung ermöglicht seit den 1930er Jahren die **Elektronenmikroskope**. Wie (b) zeigt, wird der erzeugte Elektronenstrahl zunächst mit dem Kondensor (als Quasilinse) gut gebündelt. Objektiv und Okular sind durch analoge Linsen für den Elektronenstrahl ersetzt. Alle drei elektronischen Linsen existieren magnetisch (c) oder elektrisch (d). Durch sie wird der vom Objekt beeinflusste Elektronenstrahl auf einem Leuchtschirm sichtbar gemacht. Auch die entsprechende Belichtung von chemischem, fotografischem Material ist möglich. Beim Elektronenmikroskop wirkt aber nicht die Größe der Elektronen begrenzend, sondern ihre Wellenlänge. In Abhängigkeit von der Spannung U , welche die Geschwindigkeit v bewirkt, der Elektronenmasse m_0 und der Planck-Konstante h gilt

$$\lambda = \frac{c}{f} = \frac{h}{m \cdot v} = \frac{h}{m_0 \cdot 5,93 \cdot 10^7 \sqrt{U \text{ in Volt}}}$$

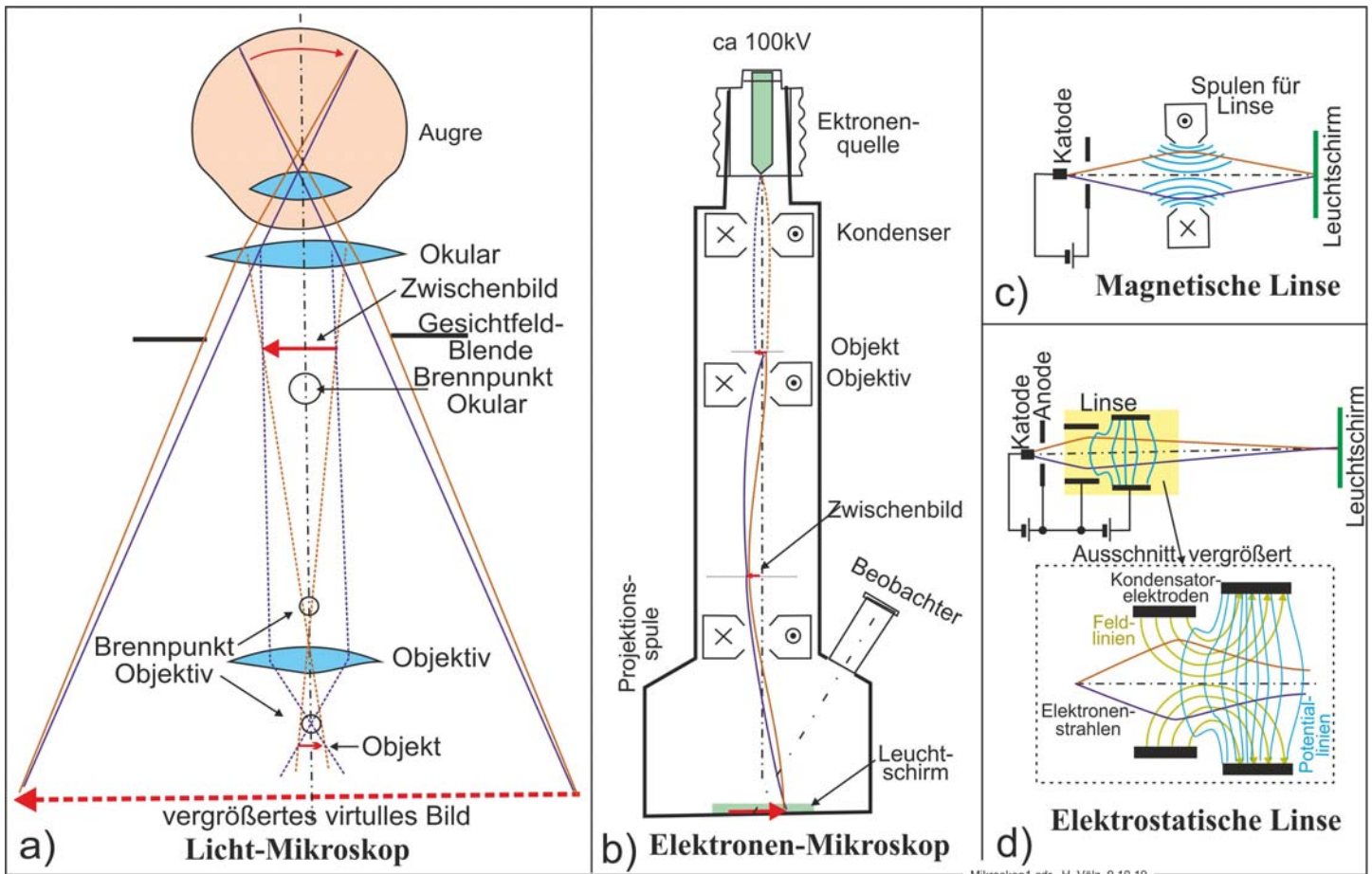


Bild 70. Lichtmikroskop (a) und Elektronenmikroskop (b) sowie Grundlagen der magnetischen (c) und elektrischen Linsen (d).

Es gibt noch mehrere elektronischen Varianten. Sie zeigt **Bild 71** im Vergleich zum Strahlengang des üblichen Verfahrens (a). Das **Feld-Elektronen-Mikroskop** (b) besitzt eine extrem feine Spitze. Auf ihr wird das sehr kleine Objekt gelagert. Von ihm gehen dann durch eine sehr hohe Spannung mittels Feldelektronenemission Elektronen geradlinig zum Leuchtschirm. Sie zeigen dann dort stark vergrößert die elektronische Oberfläche des Objektes.

Eine weitere Variante ist das **Raster-Elektronen-Mikroskop**. Ein sehr stark gebündelter Elektronenstrahl trifft auf die Oberfläche des Objektes und erzeugt dort dadurch Sekundärelektronen, die zum Empfänger gelangen. Durch die Ablenkquelle wird Oberfläche des Objektes abgerastert und parallel dazu der Elektronenstrahl der Bildröhre geführt. So entsteht nacheinander ein Oberflächenbild des Objektes. Zu diesem Verfahren gibt es Vorgänger und mechanische Sondervarianten. Bei ihnen wird die Oberfläche des Objektes mechanisch mit einer Sondenspitze ($\approx 1\text{nm}$) abgerastert. Beim **Raster-Tunnel-Mikroskop** (RTM) wird das über den Tunnelstrom ($1\text{pA} - 10\text{nA}$ bei einigen mV) gesteuert und dabei die Bewegung der Spitze registriert (d). Das **Raster-Kraft-Mikroskop** (STM) erfasst dagegen über ein Piezoelement die Bewegung der Spitze (e). Beide Varianten wurden auch für Methoden der Datenspeicherung erprobt.

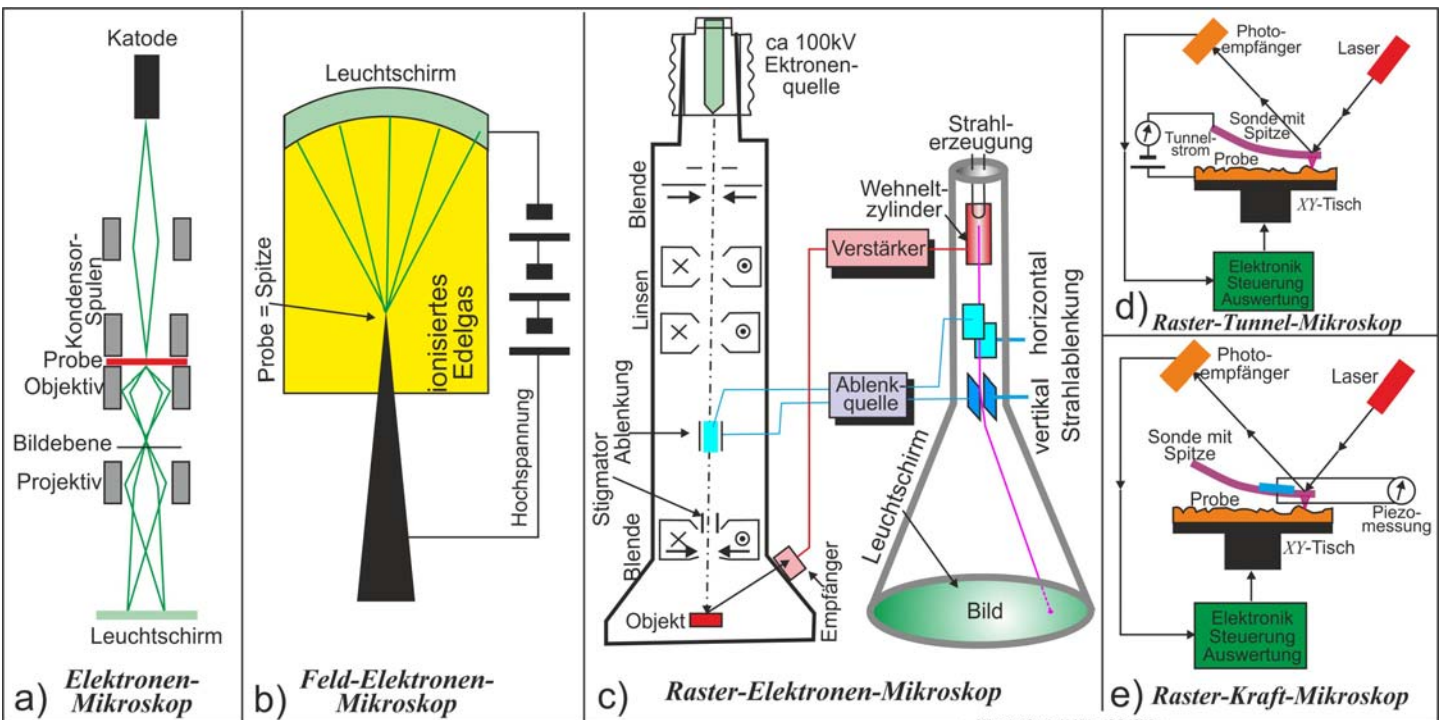


Bild 71. Weitere Varianten von Elektronen-Mikroskopen.

Außer den Licht- und Elektronenstrahl-Mikroskopen gibt es noch Varianten, die mit anderen physikalischen Verfahren arbeiten. Genauer wird noch auf das Röntgen-Mikroskop, die Computertomografie (CT), die Magnet-Resonanz-Tomografie (MRT $\approx$ Kernspintomografie) und mehr im Sinne der Fernrohre auf das Radar bei den Mikrowellen eingegangen. Weitere Methoden werden beim Schall als Sonografie behandelt. Nur erwähnt seinen dann noch das Focused-Ion-Beam-Mikroskop (FIB), Helium-Ionen-Mikroskop, Neutronen-Mikroskop, die für Wärmebilder wichtige Selbststrahlung im Infraroten, auch mit Umwandlung durch Elektronen-Vervielfacher und die erst kürzlich verfügbaren THz-Schwingungen u. a. bei Ganzkörper-Scannern.

Röntgenstrahlung entsteht, wenn schnelle Elektronen auf Materie treffen. In einem meist kleinen Brennfleck erfolgt eine starke Abbremsung, wodurch eine Strahlung mit einer Wellenlänge von 10^{-8} bis 10^{-12} m (10 nm bis 1 pm) entsteht. Bis etwa 1925 wurde die Strahlung mittels einer Gasentladung bei einem Gasdruck von 0,05 bis 0,5 bar erzeugt (a). Dabei trafen positive Ionen auf die Katode und die Röntgenstrahlung tritt durch eine Bohrung aus. Heute werden stark gebündelte Elektronenstrahlen auf eine rotierende Wolfram-Anode gestrahlt (a, b). Die Betriebsspannung reicht von 50 bis 400 kV. Die Anode erhitzt sich am Brennfleck (0,1 bis 1,5 mm) auf ≈ 2500 °C. Deshalb ist Rotationskühlung notwendig. Die Röntgenstrahlung tritt dann vorwiegend durch ein Beryllium-Fenster, selten durch Glas aus. Von der eingesetzten Energie wird nur etwa 1% in Röntgenstrahlung umgesetzt. Da die Röntgenstrahlung sich immer geradlinig ausbreitet und ist nicht ablenkbar ist erfolgen alle Abbildungen erfolgen im Durchlicht (c).

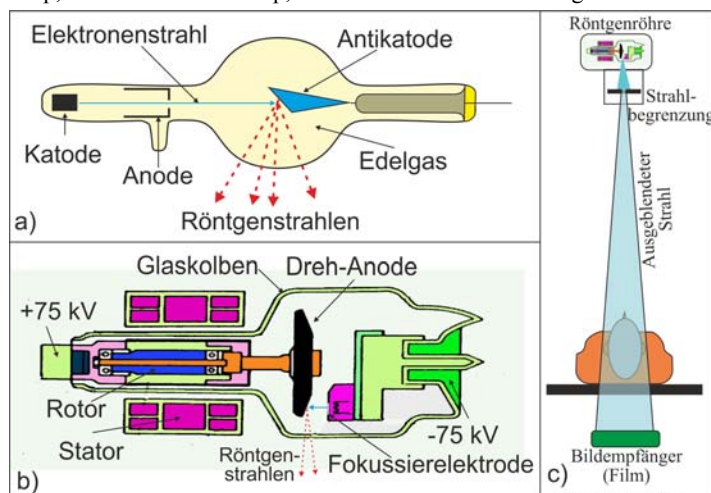


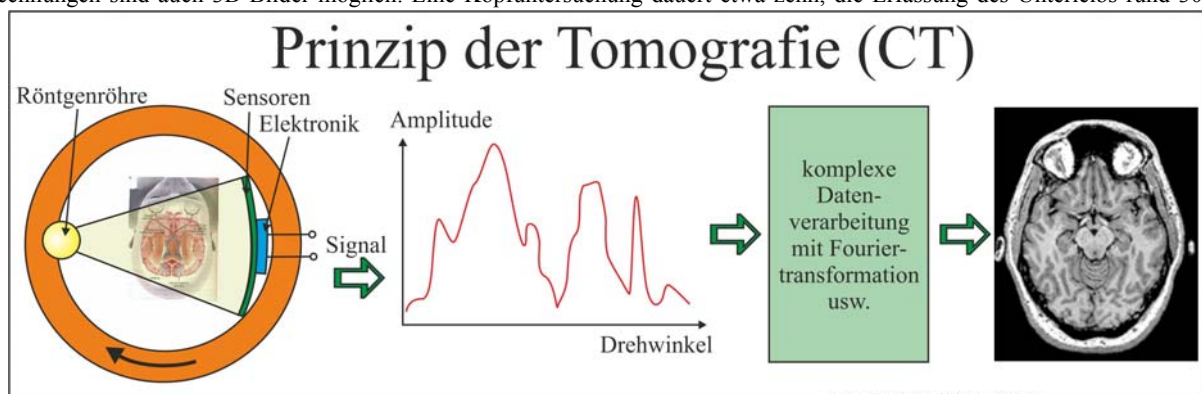
Bild 72. Das Röntgen-Mikroskop.

Tomografie

Die erste Tomografie²⁸ gab es um 1960. Die **Computertomografie** (CT) wurde aber erst 1972 von dem amerikanischen Physiker Allan M. Cormack und dem britischen Ingenieur Godfrey N. Hounsfield entwickelt. Sie erhielten dafür im Jahre 1979 den Nobelpreis für Medizin. 1972 gab es ein erstes CT-Gerät und 1974 erstes kommerzielles CT-System mit Sofortbildrekonstruktion. Mittels Röntgenstrahlen wird die örtlich unterschiedliche Durchlässigkeit der Gewebearten bestimmt. Ein fächerförmig gebündelter Röntgenstrahl wird dabei in kleinen Schritten um den Körper herum bewegt (Bild 73). Der Patient wird also nicht aus nur einer Richtung durchstrahlt, sondern nacheinander komplett rundherum aus allen Richtungen schichtweise abgetastet. Gegenüber der Röntgenquelle sind dazu auf einem Teilkreis Detektoren angebracht, welche die ankommende Restenergie messen. Zusätzlich wird der Körperteil in der runden Öffnung (Röhre) des Computertomografen längs vorgeschoben. Mit beachtlichen mathematischen Aufwand wird erst danach aus allen Aufnahmen die Rekonstruktion der Gewebeverteilung in 1 - 8 mm dicken Schichten berechnet. Mit zusätzlichen Computerberechnungen sind auch 3D-Bilder möglich. Eine Kopfuntersuchung dauert etwa zehn, die Erfassung des Unterleibs rund 30

Minuten. Zur besseren Abgrenzbarkeit ausgewählter Strukturen (Gefäße, Darm etc.) muss bei vielen Untersuchungen zu Beginn ein jodhaltiges Kontrastmittel in eine Vene gespritzt werden.

Bild 73. Prinzip des Computertomografen.



Das **MRT** (Magnet-Resonanz-Tomografie, Bild 74) wird auch MRI (magnet resonance imaging) oder NMR (nuclear magnetic resonance Spektroskopie), bzw. Kern-Resonanz, [para] magnetische Kernresonanz oder **Kernspin**-Tomografie genannt. Sie wurde im September 1971 von Paul C. Lauterbur erfunden und ist seit den 1980er Jahren im klinischen Einsatz. Sie benutzt keine Röntgenstrahlung oder Radioaktivität, aber eine enge Röhre. Sie benötigt ein sehr starkes Magnetfeld, das die Spins (Eigendrehmomente) von Wasserstoffkernen, Protonen usw. im Gewebe ausrichtet (a und b). Dabei wird eine Kreiselbewegung (Präzession) des Spins hervorgerufen (c). Die dazugehörige Larmorfrequenz hängt vom Magnetfeld ab und beträgt z. B. 42 MHz bei 1 Tesla. Ein zusätzlich eingestrahelter kurzer Impuls bewirkt dann ein Schlingern, Umklappen. Bei der Rückkehr in den Ruhezustand (Relaxation) senden die Protonen ein HF-Signal aus. Es wird bzgl. Intensität und Verzögerung gemessen. Ein zusätzliches geringeres Gradienten-Magnetfeld (d) ermöglicht die räumliche Zuordnung der empfangenen Signale. Insgesamt müssen die Magnetfelder jedoch recht kompliziert verändert werden. (klopfende Gräusche).

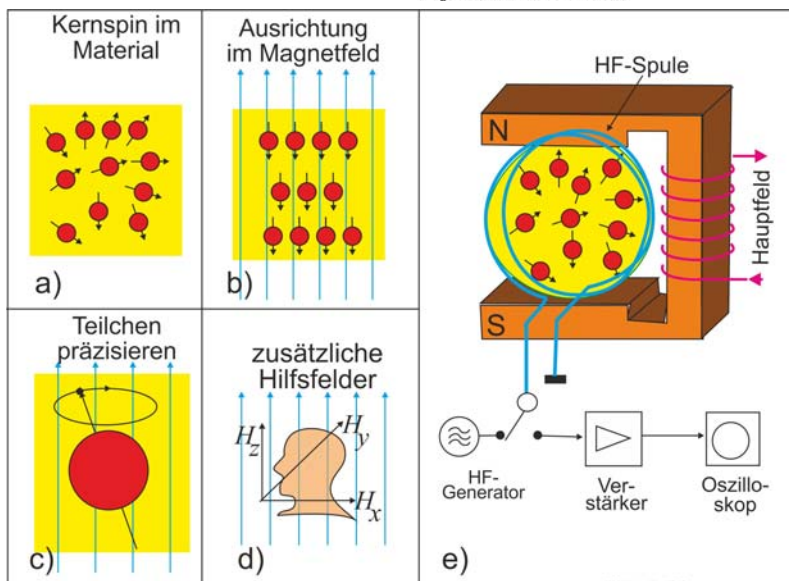


Bild 74. Zu den Grundlagen des MRT.

²⁸ Tomografie griechisch tomo Schnitt und graphein schreiben etwa synonym mit Schnittbild-oder Schichtaufnahmeverfahren

Aus allen Messwerten kann vor allem die gewebeabhängige Dichteverteilung berechnet werden. Eine Differenzierung der Gewebe kann durch den Einsatz spezieller Kontrastmittel (v. a. Gadolinium) verbessert werden. Die notwendigen hohen Magnetfelder ≥ 1 bis zu 10 Tesla fordern supraleitende Spulen und eine starke Abschirmung in der Umgebung. Was u. a. ein Entfernen von Geld, EC-Karten, Schlüssel, Metallimplantate, Herzschrittmacher usw. voraussetzt. Mit stark wachsender Leistungsfähigkeit wurde die Geschwindigkeit der Messungen erheblich gesteigert. So ist heute auch eine **funktionelle Tomografie** (fMRT) möglich. Sie gestattet vielfältige Veränderungen, z. B. Abläufe im Gehirn, bei der Durchblutung usw. zu verfolgen. Es gibt mehr Varianten des detaillierten Aufbaus und der Anwendungen.

Bei der **Positronen-Emissions-Tomographie** (PET) wird eine leicht radioaktiv „markierte“ Verbindung (Tracer) benutzt, die bei Stoffwechselprozessen vorkommt. Das typische ^{18}F -Fluorodesoxyglukose (FDG) spritzt der Arzt dem Patienten an jenen Ort (Organ), der untersucht werden soll. Alternativ können beispielsweise auch Aminosäure-Analoga verwendet werden. Der Tracer zerfällt (in Positronen) und erzeugt dabei γ -Strahlung. Dabei werden die zwei Protonen (γ -Quanten) in entgegengesetzte Richtungen abgestrahlt. Zur Analyse ist das Objekt rings von Detektoren umgeben. Jeweils zwei gegenüberliegende Detektoren empfangen dabei korreliert. Jedes so gemeinsam empfangene Signal weist den Ort zwischen ihnen und die gemessene Energie aus. Viele sich schneidende Linien ergeben schließlich ein Bild der jeweiligen örtlichen Stoffwechselrate. Daran ist u. a. zu erkennen, in welchen Regionen des Körpers ein Tumor wächst.

Funk-Antennen

In den Abschnitten zu den Teleskopen und zur Tomografie sind Radio-Wellen teilweise entscheidend. Dabei ist oft die Richtwirkung der Antennen ein wichtiger Bestandteil. In welchen Wellenbereichen dabei die verschiedenen Antennenformen bzw. Spiegel benutzt werden geht aus Bild 57, S. 29 hervor. Da es meist sehr aufwendig ist, die Richtcharakteristiken einer Sendeantenne in ihrem Umfeld messtechnisch zu bestimmen, ist das auch für Antennen geeignete Reziprozitätsgesetz nützlich. Danach ist die Richtwirkung einer Antenne beim Senden und Empfangen identisch. So wird die Richtcharakteristik meist per Empfang gemäß **Bild 75** bestimmt.

Bild 75. Bestimmung der Richtwirkung einer Sendeantenne als Empfangsantenne.

Bei den Funkantennen ist im Laufe der Entwicklung eine große Vielfalt entstanden. Sie weist die obere Hälfte von **Bild 76** aus. Ihre unterschiedlichen Richtcharakteristiken stellt vereinfacht der unter Teil dar. Dabei sind auch die erweiterten Möglichkeiten von komplex verkoppelten Antennenfeldern aufgezeigt. Ein für die Astronomie aufgebautes sehr großes Feld zeigt das (unnummerierte) Bild auf Seite 33. Eine ähnliche Zusammenschaltung mit Spiegelteleskopen ist im Bild 69, S. 34 als (e) und deren Verkopplung als (d) enthalten. Die weitaus umfangreichere Verschaltung der vielen Antennen demonstriert das Schema rechts unten in Bild 76.

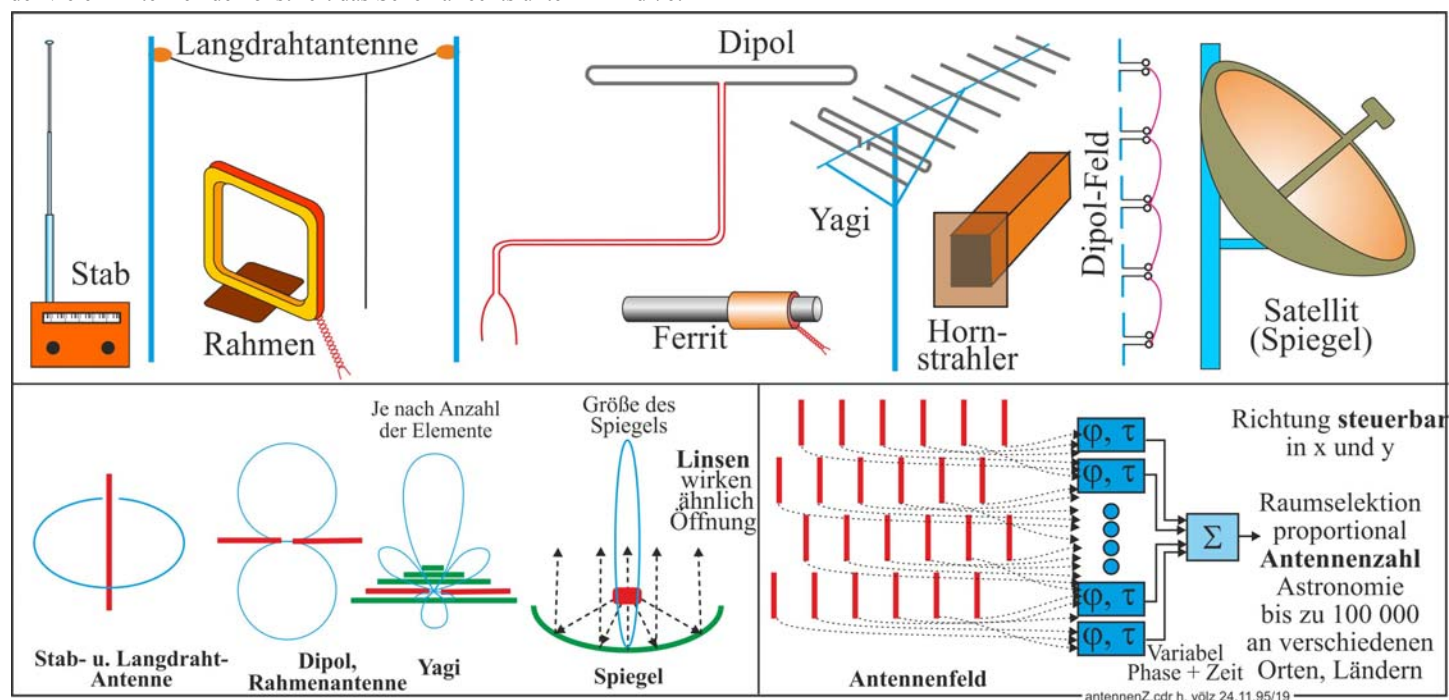


Bild 76. Vielfalt der Funkantennen und ihrer Richtungscharakteristiken sowie Verschaltung von Antennenfeldern.

Für die Satellitentechnik der Rundfunk- und Fernsehtechnik sind in Analogie und Ergänzung zu den Teleskopen von Bild 67, S. 33 weitere Varianten der Spiegelantennen entstanden. Eine Auswahl zeigt **Bild 77**. Darin ist auch eine Variante (e) mit vielen Dipolantennen enthalten.

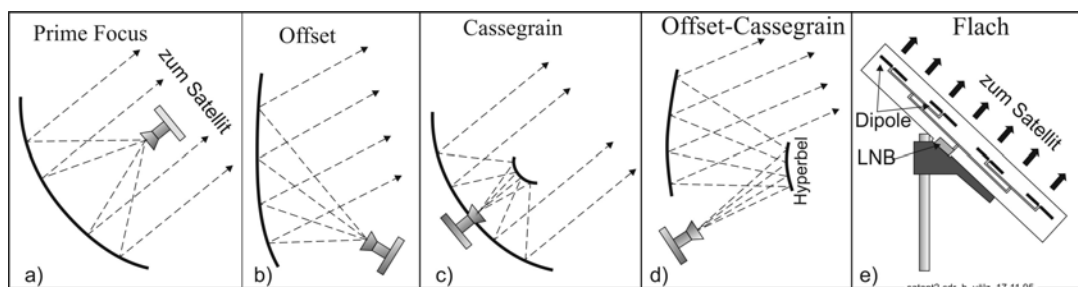


Bild 77. Bspielanordnungen von Satellitenantennen der Rundfunk- und Fernsehtechnik

Ein Beispiel für die Funk-Übertragung mit „stationären“ Satelliten zeigt **Bild 78**. Der Inhalt wird (von Frankreich) mit dem uplink zum Satelliten gesendet, dort empfangen, verstärkt und mit anderer Frequenz zur Erde zurück gesendet. Dort überstreicht er je nach Empfangspiegelgröße einen beachtlichen Bereich in Europa.

Eine deutlich andere Übertragungstechnik ist das **Scattering** von **Bild 79**. Sender und Empfänger besitzen keine direkte gegenseitige Übertragung, aber den gemeinsamen Bereich im Weltraum (lila). In der dortigen Troposphäre (≈ 10 km Höhe) kann durch Gewitter, Meteoriten, Verunreinigungen oder Flugzeuge eine Ionisierung auftreten. Über sie ist dann eine mittelbare Verbindung zwischen Sender und Empfänger möglich.

Der vielfältige Einsatz der Funkantennen – insbesondere für Handys, Smartphones usw. – verändert seit geraumer Zeit das Stadtbild. Zwei Beispiele hierzu zeigen die beiden Ausschnitte von **Bild 80**.

Bild 78. Die Funk-Übertragung per stationären Satelliten.

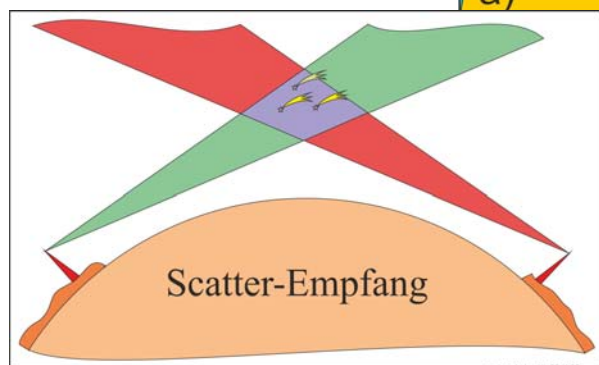
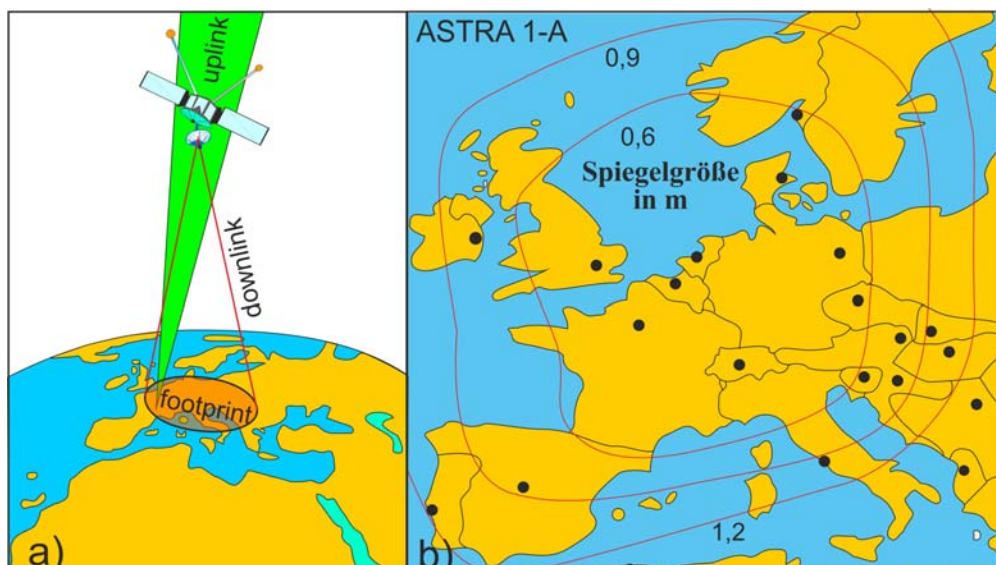


Bild 79. Prinzip der Scattering-Übertragung.

Eine gewisse Ähnlichkeit zu den Tomografien ab S. 36 besitzt das **Radar** (Radio Detection and Ranging). Früher wurde es auch Radio Aircraft Detection and Ranging genannt. Anfangs war die deutsche Bezeichnung „Funkmeßtechnik“ üblich. Erste Versuche zur Ortung mit Radiowellen führte Christian Hülsmeyer bereits 1904 durch. Sie gerieten aber in Vergessenheit. So ist der schottische Physiker Sir Robert Alexander Watson-Watt (1892–1973) eigentlicher Erfinder. Er erforschte damit Reflexion von Radiowellen in der Meteorologie (1919 Patent). 1935 erfolgte erstmals die Radar-Ortung von Flugzeugen im Meterwellenbereich. Dafür war die Erfindung des Magnetrons an der Universität Birmingham Anfang 1940 ganz wesentlich. Heute werden hauptsächlich unterschieden: Das **Impulsradar** mit kurzen Impulsen konstanter Frequenz und dazwischen der Empfang des Echos über die gleiche Antenne und das **Dauerstrichradar** CW (continuous wave) mit Frequenzmodulation und konstanter Amplitude. Beim letzten wird der Doppler-Effekt zur Geschwindigkeitsmessung ausgenutzt. Das Prinzip des Impulsradars zeigt **Bild 81**. Der Sender erzeugt die auszustrahlenden Schwingungsimpulse. Sie werden mit unterschiedlicher Verzögerung von den einzelnen Objekten reflektiert und nach dem Empfang verstärkt. Da die Parabelspiegel-Antenne rotiert, ermöglichen die rückkehrenden Impulse ein Bild der Umgebung. Dabei können für die dort vorhandenen Objekte folgende Aussagen gewonnen werden: Raumwinkel bzw. Richtung, Entfernung und Relativbewegung und deren Geschwindigkeit mittels Doppler-Effekt. Bei guter Auflösung können aus den Konturen z. B. der Flugzeugtyp usw. abgeleitet werden. Deshalb müssen Radarantennen eine stark ausgeprägte Richtungscharakteristik aufweisen. Die Breite des Strahles ist proportional der Wellenlänge und umgekehrt proportional zur Breite der Antenne. Zur Vermeidung großer Antennen bei mobilen Geräten wird daher ein Mikrowellenradar eingesetzt. Der Dopplereffekt wird analog bei Geschwindigkeitsmessung im Autoverkehr sowie in der Natur mittels Schall u. a. bei Fledermäusen und Walen genutzt. Die Frequenzbereiche der Funktechnik konnten nur schrittweise in die Nutzung einbezogen werden. Hierzu gibt **Bild 82** einen Überblick.

Bild 81. Prinzip des Impuls-Radars (rechts oben).

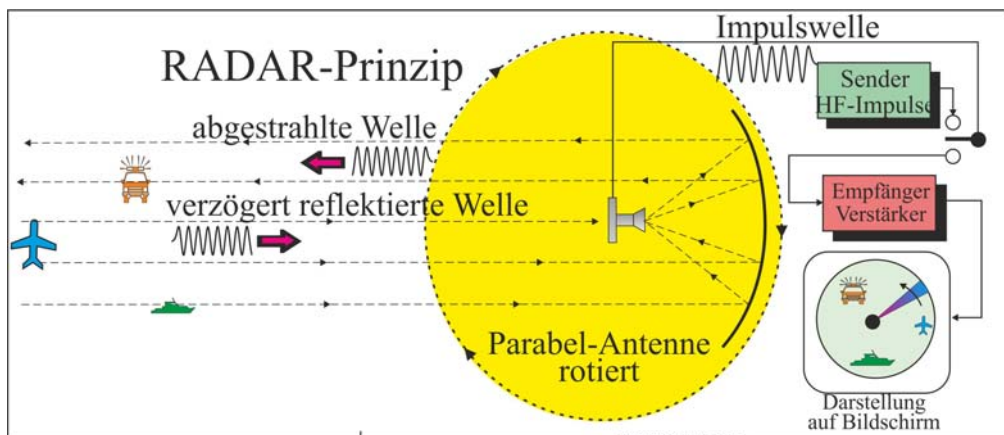


Bild 80 Typische Antennenwälder auf den Dächern der Stadt.

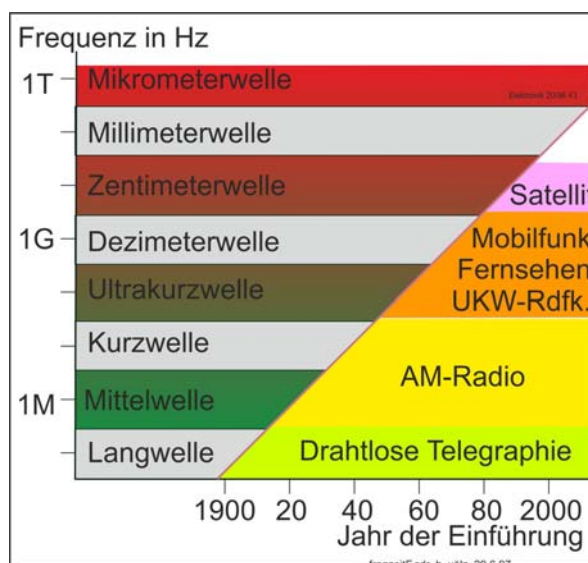


Bild 82. Erschließung der Frequenzbereiche der Funktechnik.

Möglichkeiten bei Schall

Wird von recht kurzen Schallsignalen abgesehen, so ist Schall immer nur als fortlaufendes Geschehen inhaltlich verständlich, insbesondere als Sprache oder Musik. Ansonsten ist er ein vielfach nicht interpretierbares Geräusch. Jedoch sein Nachhall oder Echo kann zusätzliche Aussagen zum Raum und dort Vorhandenes liefern. Das benutzen z. B. Fledermäuse und Wale. Es ermöglicht auch analoge technische Anwendungen, die unten u. a. als Sonografie noch behandelt werden. Mittelbar wurde der Schall auch sehr früher intuitiv für „Ungewöhnliches“, z. B. für religiöse Zwecke benutzt. Über Kanäle, Rohre usw. konnte so scheinbar ein Heiligtum oder eine Wand sprechen und sogar Fragen beantworten. Recht früh wurden auch zufällig Flüsterecken, -bögen usw. entdeckt. Delphi war der Sitz des Orakels der Erdgöttin Gaia. Das Zentrum war die oberste Priesterin Pythia, deren ekstatische Weissagungen, von den Priestern in Verse übersetzt wurden. Zu den Weissagungen kamen sogar Herrscher nach Delphi. In Kreta gibt es die Grotte „Ohr des Dionysos“, die teilweise zur Erzgewinnung aus dem Gestein bearbeitet wurde. Ihre ungewöhnliche Akustik wird immer noch von Reisgruppen ausprobiert. Angeblich hat hier der Diktator von Syrakus, Dionysos seine gefangenen Feinde im Steinbruch schreien lassen und ihnen die Grotte nachts als Unterkunft gewährt. Aufgrund der besonderen Akustik soll er so abends ihre geflüsterten Gespräche belauscht haben. Auf ihn bezieht sich der Dionysos in Schillers Bürgschaft.

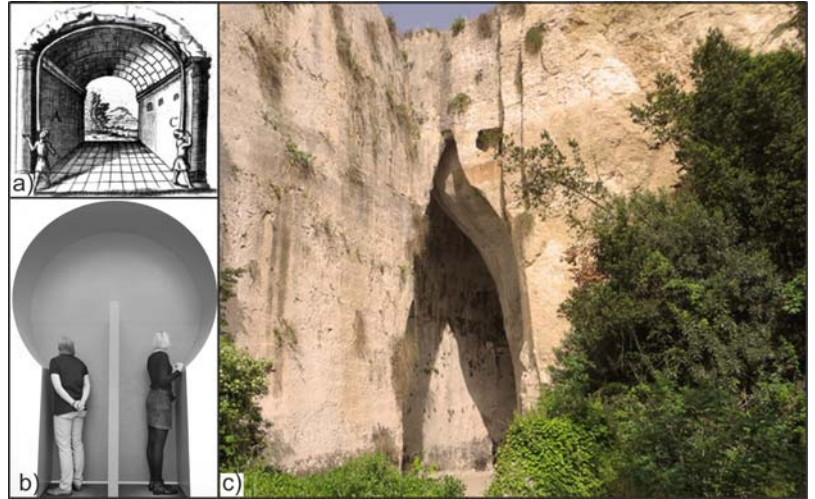
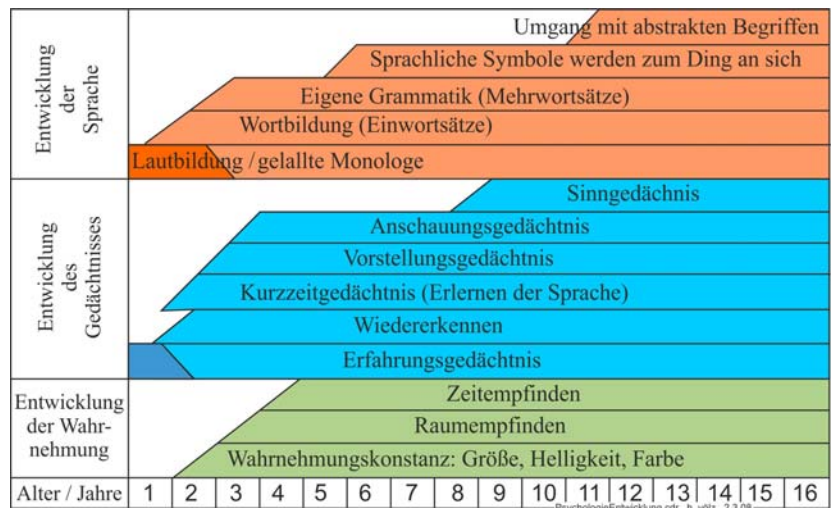


Bild 83. Lauschbögen (a, b) und Ohr des Dionysos (c)

Doch was Schall eigentlich ist, blieb sehr lange unklar. Die **Akustik**²⁹ als Lehre vom Schall entstand frühestens ab Mitte des 17. Jh. Wahrscheinlich angeregt durch Gewitter schlug Francis Bacon (1561 - 1626) vor, den zeitlichen Abstand von Blitz und den Knall bei Kanonen zur Bestimmung der Schallgeschwindigkeit zu nutzen. Das gelingt – noch über Messung mit dem Pulsschlag Marin Mersenne (1588 - 1684) und Pierre Gasendi (1592 - 1655).

Obwohl ein Kind bereits im Mutterleib Schall recht gut und sogar als einzigen Außenreiz wahrnimmt, muss es ganz im Gegensatz zur Evolution sofort und primär die Gesamtheit des Schalls hören. Einzelheiten wie Tonhöhe, Takt, Rhythmus, Nachhallzeit, Lautstärke usw. müssen erst erlernt werden (Bild 84). Dennoch soll u. a. Musikalität dann deutlich gefördert werden, wenn ein Kind bereits im Mutterleib viel entsprechende Musik hört. Insgesamt ist es daher besonders herausragend, dass bereits Pythagoras (um 570 bis circa 500 v. Chr.) den Zusammenhang zwischen der Saiten- bzw. Pfeifenlänge und den musikalischen Intervallen erkannte.

Bild 84. Kindliche Entwicklung zur Schallwahrnehmung.



Für eine unmittelbare Anwendung (Verstärkungen) von Schall existieren nur wenige Beispiele. Fast immer sind elektronische Methoden mit Mikrofonen, Lautsprechern oder Kopfhörern notwendig. Recht oft erfolgt auch eine zusätzliche Umsetzung in Bilder. Außerdem werden Frequenzen außerhalb unseres Hörbereiches, also Infra- und Ultraschall genutzt. Allgemein kommen auch beim Schall Methoden zur Anwendung, die analog zu den Linsen, Spiegeln und Antennenfeldern der Optik und Nachrichtentechnik sind (s. u.). Dabei ist aber zu beachten, dass beim Schall eine longitudinale statt der elektromagnetisch transversalen Welle vorliegt (vgl. Bild 58, S. 29). Außerdem treten als Folge der viel geringeren Schallgeschwindigkeit deutlich größere Wellenlängen auf (vgl. Bild 57 S. 29). Weiter ist zu beachten, dass die Absorptionskoeffizienten für Schall meist viel geringer als für Licht sind. Dadurch erfolgen zumindest in Räumen umfangreicher Reflexionen und damit komplizierte Schallfelder.

Fast alle elektro-akustischen Wandler funktionieren nach dem gleichem Prinzip: Eine Membran oder Oberfläche hat direkten Kontakt zum Schallfeld und ihre mechanischen Bewegungen sind mit dem elektrischen Signal verknüpft. Auch hierbei gilt wieder – von wenigen Ausnahmen abgesehen – die Reziprozität. Daher kann fast jeder konkrete Wandler für alle drei o. g. Anwendungen benutzt werden und hat dann als Lautsprecher, Kopfhörer oder Mikrofon gleichen Frequenzgang und gleiche Richtungsabhängigkeit. Infolge der beachtlichen Leistungsunterschiede, insbesondere zwischen Mikrofon und Lautsprecher sind jedoch meist umfangreiche angepasste Dimensionierungen und Zusatzmaßnahmen erforderlich, die sich dann auf Frequenzgang und Richtungsabhängigkeit dann auswirken können. Für die Verkopplung zwischen den Membranbewegungen und elektrischen Signalen sind mehrere Aufbauprinzipien entstanden. Da sie bei Lautsprechern besonders vielfältig sind, seien zunächst ihre Varianten betrachtet. Auf die Mikrofone und Kopfhörer wird anschließend eingegangen.

In Bild 85 sind die verschiedenen Methoden zur Schallerzeugung zusammengestellt. Auch die historischen Varianten sind dabei einbezogen. So wird heute das **adiabatische Ionenprinzip** (a) kaum noch angewendet. Es ist auch unter Namen „singender Lichtbogen“ bekannt. Bei ihm tritt primär keine Bewegung auf. Ein Hochfrequenzgenerator HF erzeugt einen Lichtbogen. Dadurch erwärmt sich die Luft in seiner unmittelbaren Umgebung und dehnt sich aus. Wird die HF-Amplitude mit dem Schallsignal moduliert, so ändern sich die Temperatur des Lichtbogens und damit auch die Ausdehnung der umgebenden Luft. Rein theoretisch ist dieses quasi masselose und damit resonanzfreie Prinzip ein idealer Lautsprecher, denn es besitzt keinen Frequenzgang. Leider sind aber der Wirkungsgrad sehr schlecht, die erzielbare Lautstärke gering und der Aufwand sehr hoch. Dennoch wird es zuweilen bei sehr hohen Frequenzen benutzt. Auch das zweite historische Prinzip mit der **elektrostatische Haftreibung** (b) wurde fast nur in der Anfangszeit des Rundfunks genutzt. Es beruht auf den Johnson-Rahbeck-Effekt der gleitenden Reibung. Mit der Spannung ändert sich Reibungskoeffizient zwischen der sich drehenden Achatwalze und dem federgespannten Stahlband. So wird das Band gegenüber der Federkraft bewegt und beeinflusst die Membran. Früher wurde auch das magnetische **Freischwingerprinzip** (c und d) häufig angewendet. Ein Anker, auch Zunge genannt, befindet sich in einem Dauermagnetkreis. Der Strom durch eine Spule verändert das auf den Anker wirkende Magnetfeld, wodurch sich der Anker seitlich bewegt. Die Ankopplung zur Membran, meist ein Trichter, erfolgt an dem im Bild mit einem Pfeil angedeuteten Punkt. Durch einen

²⁹ griechisch akouein hören

anderen Aufbau erfolgt u. a. eine unmittelbare Anregung der Membran (e). Diese Variante wird heute zuweilen noch ohne Trichter bei einfachen Kopfhörern benutzt. Eine weitere Abwandlung ersetzt die Spule durch einen einzigen **Leiter im Feld** des Dauermagneten (g). für die Schallabstrahlung oder -aufnahme ist der Leiter in ein gefaltetes und leicht bewegliches Bändchen umgeformt. Nach diesem Prinzip funktionierten die ersten hochqualitativen **Bändchen-Mikrofone** (Details s. Bild ##, S.##.) Eine beachtliche Abwandlung hiervon zeigt (l). Das Bändchen ist dabei zur **Spule auf einer Membran** abgewandelt. Dieses Prinzip wird heute teilweise bei Kopfhörern und Hochtönlautsprechern genutzt. Eine erhebliche Modifizierung der magnetischen Wirkung erfolgt mit dem **Tauchspulenprinzip** (f). Es wird heute fast ausschließlich und in beachtlichen vielen Abwandlungen bei den **elektrodynamischen Lautsprechern** benutzt (weitere Details s. u.). Schließlich ist noch die **elektrostatische Variante** zu nennen (k). Zwischen zwei gegenüberstehenden durchlöcherten Elektroden befindet sich die vom Strom durchflossene Membran. Die Elektroden werden mit der Signalspannung versorgt, wodurch wird die Membran wechselseitig stärker von den einzelnen der Elektroden angezogen wird.

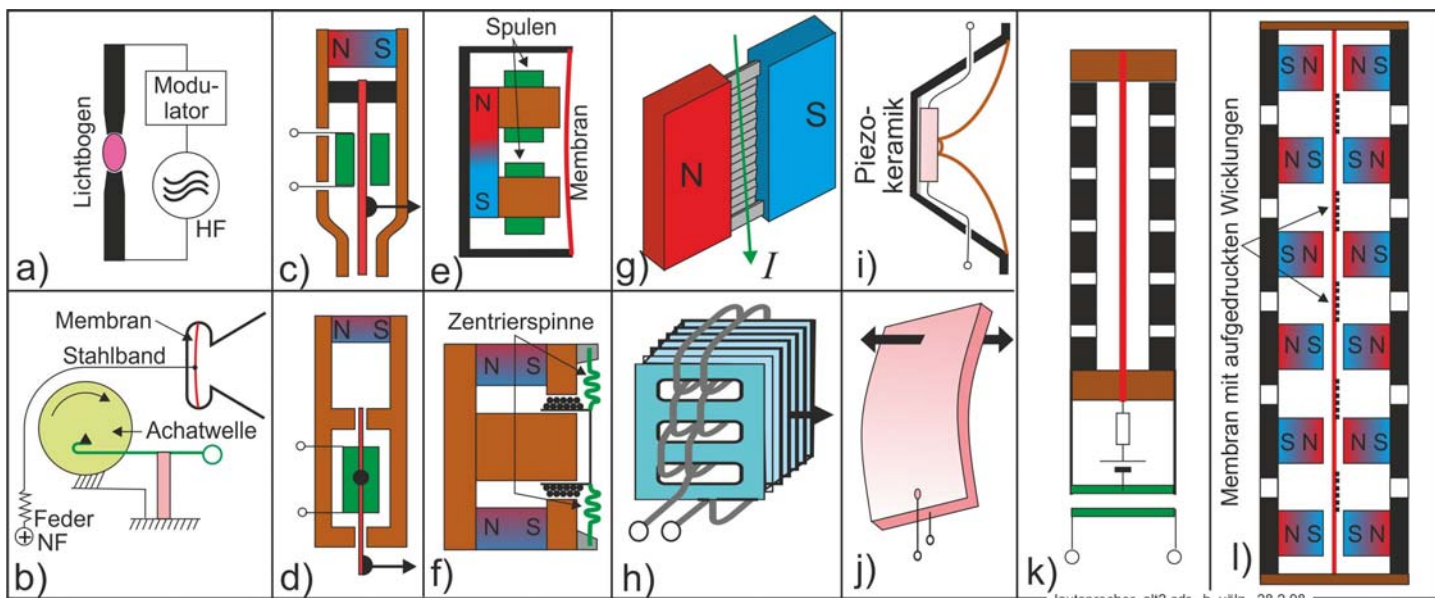


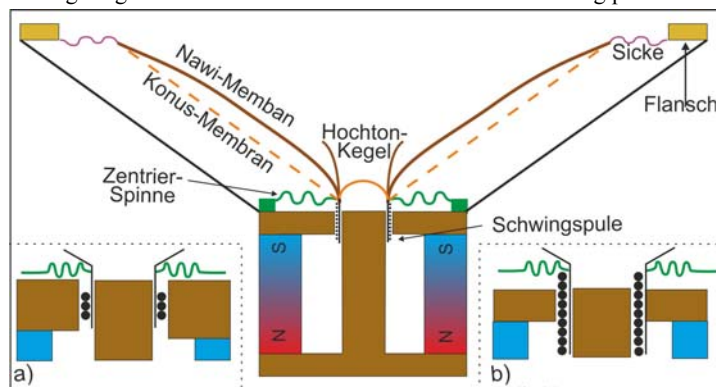
Bild 85. Die wichtigsten Prinzipien zur Erzeugung von Schall.

Der elektrodynamische Lautsprecher

Der elektrodynamische Lautsprecher wurde bereits 1898 vom Engländer Oliver Lodge erfunden. Doch zunächst erschien sein Aufbau so kompliziert und sein Wirkungsgrad so gering, dass er sich erst in den späten 30er Jahren durchsetzte. Weil ein fremderregter Magnetkreis zu hohem Energieverbrauch forderte waren die neuen Hochleistungsmagnete eine ganz wesentliche Voraussetzung. Den heute üblichen Aufbau zeigt Bild 86. Ein Ringmagnet, zunächst aus Ferrit dann AlNiCo und heute weitgehend magnetische Spezialwerkstoffe, erzeugt über Weicheisenjochs einen engen zylindrischen Spalt mit möglichst hoher Feldstärke. In ihm ist die bewegliche, ein- oder zweilagige **Schwingspule** gelagert. Sie soll sowohl einen möglichst geringen elektrischen Widerstand als auch eine kleine Masse besitzen. Allgemein wird runder Kupferdraht verwendet. Einen kleineren Widerstand und eine bessere Ausnutzung der Spaltfeldstärke ermöglicht **rechteckförmiger** aber teurer Kupferdraht. Für einen typischen Tieftönlautsprecher mit 30 cm Konusdurchmesser enthält die Schwingspule ca. 10 m Draht. Bei ihrer Bewegung wirkt die Schwingspule auch als Generator und erzeugt Strom. Er wird bei Kurzschluss besonders groß. So können Eigenschwingungen des Systems gedämpft werden. Deshalb werden ein kleiner Innenwiderstand des Verstärkers und „dicke“ Lautsprecherleitungen gefordert. Die richtige Lage der Schwingspule im Spalt wird durch eine elastische Zentrierspinne erzwungen. Die Auslenkung der Schwingspule ist solange dem eingespeisten Strom proportional, wie sie sich ganze im konstanten Spaltfeld befindet. Für den Membranhub gilt dann $h \approx N/(25 \cdot A \cdot f)$ darin ist A die Membranfläche, f die Frequenz und N die Leistung. Für einen schon großen Lautsprecher mit 40 cm Membrandurchmesser erzeugen bereits nur 0,1 W bei 30 Hz einen Schwinghub um 2 mm. Wenn dann die Schwingspule etwa genauso lang wie der Luftspalt ist, müssen ihn zeitweilig einige Windungen verlassen. Das führt zu deutlichen Nichtlinearitäten. Daher wird die Schwingspule deutlich abweichend vom Luftspalt gewählt (a, b). Dann bleiben auch bei größeren Hüben immer gleich viele Windungen im Luftspalt. Beide Varianten senken den Wirkungsgrad, kurze Spule das Magnetfeld, lange den wirksamen Spulenstrom.

Die Schwingspule treibt die **Membran** an. Dazu soll sie leicht und steif sein. Deshalb kommen langfasriges Papier, Kunststoff und Hart-schaum (Polysterol) sowie Titan- oder Aluminiumlegierungen (z. T. in Wabenstruktur) zur Anwendung. Ursprünglich hatte die Membran die Form eines Kegelausschnittes (Konusmembran). Ein großer Öffnungswinkel macht dabei die Membran leicht und etwas steif, ein kleiner bewirkt eine höhere Masse. Es treten auch leicht störende Teilschwingungen innerhalb der Membran auf. Deshalb ist Polypropylen mit seiner größeren inneren Dämpfung vorteilhaft, obwohl es schwerer als Papier ist. Auch eine Papiermembran beschichtet mit geeignetem Kunststoff ist vorteilhaft. In jedem Fall bringt eine Nawi-Membran (**n**icht **a**bwickelbar) erheblichen Schutz gegenüber Eigenschwingungen. Dabei ist sie in der Nähe der Schwingspule relativ dick (steif) und zum Außenrand hin verdünnt. So schwingt sie bei tiefen Frequenzen vollständig homogen, bei hohen Frequenzen dagegen nur in der Nähe der Schwingspule. Eigentlich wäre zwar eine plane Membran für den Übergang der Schallschwingungen zur Luft ideal, weil sie eine ebene Schallfront erzeugt. Außerdem ergäbe sich dann ein Flachlautsprecher mit geringer Einbautiefe. Deshalb wurde sie auch zeitweilig produziert. Zur Erhöhung ihrer Steife erhielt sie dabei eine komplexe Wabenstruktur. Dennoch war keine ausreichende Unterdrückung der Membranverbiegungen zu erreichen. Das Taumeln der Schwingspule wird durch die zweite „Führung“ mit der **Sicke** am äußeren Rand weitgehend verhindert. In einigen Fällen wird zusätzlich ein besonders steifer **Hochtönkegel** angebracht. Insbesondere bei den leistungsstarken PA-Lautsprechern (public address – große Übertragungen) wird der gesamte Spalt mit einem **Ferrofluid** ausgefüllt. Es verbessert die Wärmeableitung. Außerdem vergrößert seine magnetische Komponente die wirksame Spaltfeldstärke. Schließlich wird das ganze Lautsprecherchassis durch einen stabilen metallischen **Korb** gehalten. Wegen der Membran wird er auch **Konuslautsprecher** genannt.

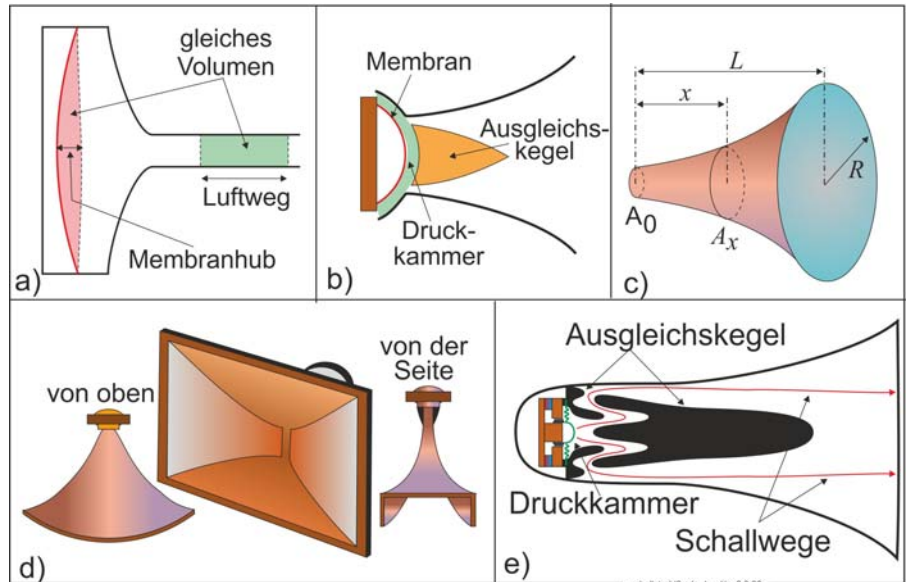
Bild 86. Aufbau des elektrodynamischen Lautsprechers.



Schallabstrahlung und Frequenzgang

Ein Nachteil aller Konuslautsprecher ist der schlechte Wirkungsgrad. Es werden praktisch nur 3 bis 5 % erreicht. Die Ursache liegt vor allem darin begründet, dass die Membran sehr dem Schallwellenwiderstand der Luft sehr schlecht angepasst ist. Als Folge weicht die Luft weitgehend seitlich aus. Das war ein wesentlicher Grund für die späte Verbreitung der Konuslautsprecher. Mittels Drucktransformation und Exponentialtrichter lässt sich jedoch der Wirkungsgrad deutlich erhöhen. Die **Druckkammer** nutzt die Kontinuitätsgleichung (**Bild 87**). Die Membranbewegung erfolgt mit kleiner Amplitude und großer Fläche (a). Dann wird der Schall in einen engeren Querschnitt gezwungen, wobei seine Amplitude nahezu proportional mit der Flächenverringernung zunimmt. Die gebräuchliche technische Ausführung zeigt schematisch (b). Mit dem **Exponentialtrichter** (c) wird dann die primäre Schwingung dem Wellenwiderstand der Luft durch eine kontinuierliche Transformation angepasst. Mit der Trichterkonstanten k für den Querschnitt am Ort x gilt dann $A_x = A_0 \cdot e^{kx}$. Hierin ist A_0 der Querschnitt der Anfangsöffnung (Halsöffnung) des Trichters. Die tiefe übertragbare Frequenz ist durch die Endöffnung (Mund) bestimmt. Für 30 Hz werden 5 m Weite und damit etwa 10 m^2 Querschnitt benötigt. Solche gewaltigen Trichter waren in der Anfangszeit des Kinos durchaus gebräuchlich. Durch die erheblich bessere Anpassung wird jedoch der Wirkungsgrad etwa verzehnfacht. Zusätzlich tritt der Schall aus dem Exponentialtrichter vorteilhaft als ebene Wellenfront aus. Deshalb war der Exponentialtrichter auch beim Grammophon so wichtig. Heute werden Trichter nur noch für Hoch- und z. T. für Mitteltöner (s. u.) eingesetzt. Dann liegen die Endöffnung im Dezimeterbereich und die untere Grenzfrequenz bei mehreren hundert Hz. Hier sind dann leicht verzerrungsarm Schalldrücke von 100 bis 120 dB zu erreichen. Neben den kreisförmigen Trichtern sind auch leichter herstellbare rechteckige entstanden (d), die auch Hörner genannt werden. Zur Verkürzung der Trichterlänge wird der Schallstrahl beim Druckkammerlautsprecher mehrfach umgelenkt (e). Gegenüber früher dominiert die Forderung nach höchstmöglicher Klangqualität, also nach geringsten nichtlinearen Verzerrungen und umfangreichen gleichmäßigem Frequenzbereich. Trichter werden fast nur noch bei Hochtönern (kleine Öffnung) oder für besonders große Lautstärken z. B. bei Megafonen eingesetzt.

Bild 87. Druckkammer und Exponentialtrichter.



Jeder Lautsprecher strahlt mit seiner Membran Schall gleichzeitig, aber gegenphasig nach vorne und hinten ab (**Bild 88a**). Treffen beide Schallwellen aufeinander, so verstärken oder schwächen sie sich je nach ihrer Phaselage. Das ist durch die Frequenz f , bzw. Wellenlänge λ und der Wegdifferenz Δx bestimmt. Eine Schallabstrahlung wird dann weitgehend unterdrückt, wenn $\Delta x < \lambda/4$ gilt, was besonders stark die tiefen Frequenzen betrifft. Häufig wird dabei von akustischem Kurzschluss gesprochen. Um das zu vermeiden, müssen der Vorder- und Rückschall möglichst gut voneinander getrennt werden. Einige Zusammenhänge dafür sind besonders übersichtlich bei einer Schallwand, vorwiegend aus Holz oder Kunststoff. Für einige Ausführungen zeigt (b) die Frequenzgänge. Bei einer kreisförmigen Schallwand treten ausgeprägte Maxima und Minima auf. Bereits zwei optimal voneinander abweichende Radien lassen einen deutlich besseren Frequenzgang entstehen. Noch günstigere können rechteckige Schallwände sein. Die zwei Beispiele zeigen, dass auch die Lage des Lautsprechers in der Schallwand von erheblichem Einfluss ist. Für die gerade noch brauchbar abstrahlende tiefste Frequenz muss für den kürzesten Abstand zum Rand hin etwa $r [\text{m}] \geq 70/f [\text{Hz}]$ gelten. Für 100 Hz sind also bereits reichlich 3 m^2 und für 40 Hz sogar 20 m^2 erforderlich. Daher mussten für tiefe Frequenzen andere wirksame Lösungen gefunden werden.

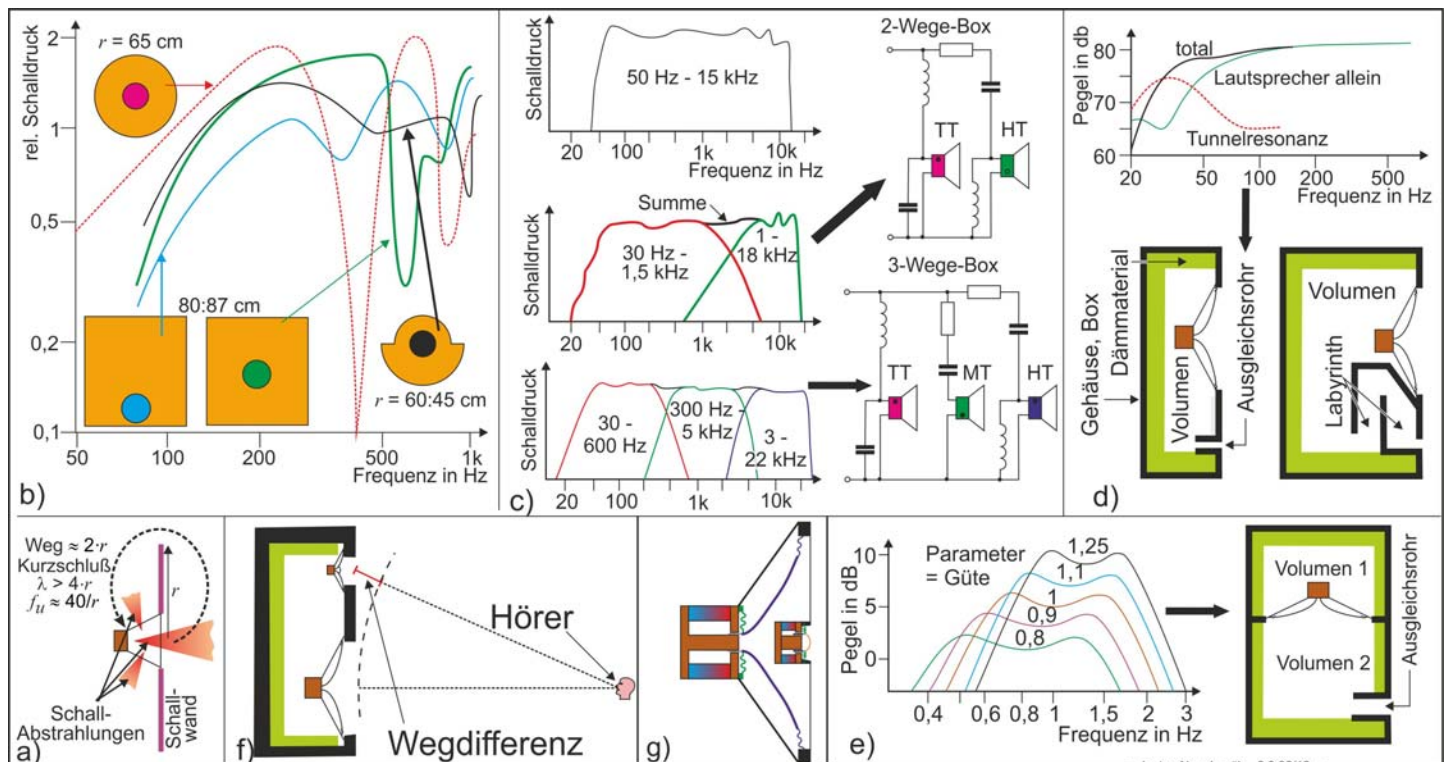


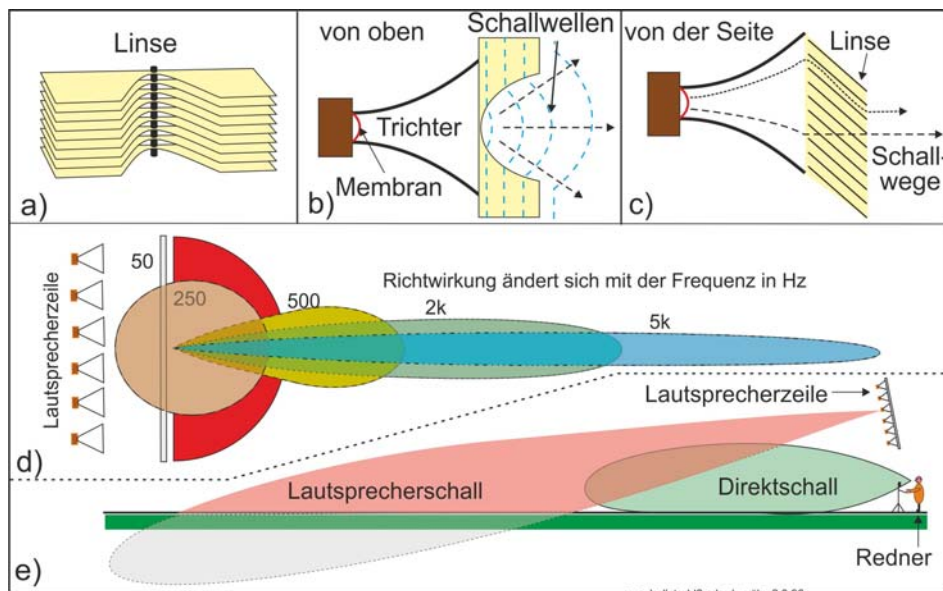
Bild 88. Maßnahmen zur Verbesserung des Frequenzganges durch Lautsprechersysteme.

Auf die Trichter wurde bereits oben hingewiesen. Heute werden aber überwiegend mehrere Lautsprecher in Teilfrequenzbereichen eingesetzt. Beispiele dafür zeigt (c). Dabei sind zu unterscheiden: Hochtonlautsprecher (auch Hochtöner oder Tweeter genannt) strahlen bei etwa 3 bis 22 kHz, Mitteltonlautsprecher (Mitteltöner oder Squawker) bei ca. 300 bis 5000 Hz und Tieftonlautsprecher (Tieftöner, Woofer oder Subwoofer) bei ca. 20 bis 600 Hz. Ihre Spannungen erhalten sie dann über Filter. Durch leicht überlappende Bereiche kann insgesamt ein leidlich geradliniger Frequenzgang erreicht werden. Da meist die tiefen Frequenzen Probleme bereiten, haben sich Boxen mit Bassreflex (Resonanz) durchgesetzt. Mit der Rohröffnung und -länge (Tunnel) sowie dem eingeschlossenem Volumen entsteht eine Resonanz bei der tiefsten zu übertragenden Frequenz (d). Zuweilen wird dabei das Rohrsystem auch durch ein Labyrinth ersetzt. Mittels der Resonanz und den Eigenschaften des Lautsprechers können deutlich tiefere Frequenzen abgestrahlt werden. In Sonderfällen wird das innere Volumen auch zweigeteilt. So lässt sich für die Abstrahlung ein Bandpassbereich erreichen (e). Bei der Benutzung mehrerer Lautsprecher tritt häufig eine störende Laufzeit auf (f). Deshalb wird zuweilen der Hochtöner vor dem Tieftöner aufgehängt (g). Trotz aller Maßnahmen und erheblicher Verbesserung der Lautsprechersysteme ist elektroakustisch die Schallabstrahlung immer noch das schwächste Glied für hohe Ansprüche.

Die Beeinflussung der **Richtwirkung** bei der Schallabstrahlung ist schwierig. Das hängt mit den im Vergleich zur Optik erheblich größeren Wellenlängen zusammen. Daher erfolgt sie auch nur bei wenigen Sonderanwendungen und vor allem bei hohen Frequenzen. Mit einem Trichter

geschieht es z. B. nur bei mittleren Frequenzen, etwa bei Megaphonen (s. o.). Prinzipiell sind über teilweise „Umwege“ Quasilinsen zu erzeugen. Ein Beispiel zeigt **Bild 89**. Einen möglichen Aufbau – der vor dem Lautsprecher, bzw. (-trichter) eingefügt wird – zeigt (a). Dabei bleibt die seitliche Abstrahlung – wie es (b) von oben gesehen zeigt – nahezu unbeeinflusst. Jedoch von der Seite gesehen (c) wird der seitlich hindurchgehende Schall (punktirt) stärker als der in Mitte (gestrichelt) verzögert. So entsteht eine beachtliche Richtwirkung. Eine andere Möglichkeit entsteht durch Lautsprecherzeilen (d). Je nach der Wellenlänge bündeln sie den Schallstrahl. Zusätzlich kann so die Luftdämpfung bei höheren Frequenzen deutlich und ihre Hörweite verlängert werden. Weiter kann bei Übertragungen im Freien die begrenzte Reichweite des Schalls deutlich verändert werden (e).

Bild 89. Beispiele für die Beeinflussung der Richtung des Schalls.



Mikrofone

Für die Schallübertragung ist ein Mikrofon erforderlich. Es benötigte 1867 Bell dringend für sein Telefon. Bereits 1858 hatte Théodor du Moncel das **Kohlemikrofon** vorgestellt. Für wenig später werden aber auch Hughes und Edison als Erfinder genannt. Beim Kohlemikrofon befinden sich zwischen der elektrisch leitenden Membran und der leitenden Gegenelektrode kleine Kohlekörner bzw. Kohlegries (**Bild 90a**). Der Schalldruck presst sie periodisch zusammen und ändert dadurch den Widerstand zwischen Membran und Gegenelektrode. Seit etwa 30 Jahren hat es nur noch historische Bedeutung. Deutlich später entstanden **elektrodynamische** Mikrophone, die eine hohe Anwendungsbreite erlangten. Ihre Urform ist das **Bändchenmikrofon** von 1906 (b). Es war das erste hochwertige Studiomikrofon. Seine „Membran“ ist nur ein dünnes geriffeltes Aluminiumbändchen, dass im Luftspalt eines Magneten leicht jeder Schallschwingung folgen kann. Die so extrem geringen Induktionsspannungen werden bereits im Innern des Mikrophons durch einen Transformator deutlich vergrößert. Auch es hat heute keine praktische Bedeutung mehr. Lange Zeit waren einfache **magnetische Mikrophone** führend. Teilweise entsprechen sie der Umkehr des Kopfhörers (s. u. bzw. Bild 85e). Sie entsprachen teilweise aber auch dem Freischwinger-Lautsprecher (c). Die Membran ist dann starr mit der beweglichen magnetisierbaren Zunge verbunden. Bei Schallschwingungen nähert sie sich abwechselnd den Leitblechen des Nord- bzw. Südpols des Magneten. Auch **Kristallmikrophone** wurden und werden zuweilen sogar noch benutzt. Durch den Schall wird eine Piezokeramik verbogen (d). Die entstehende Spannung entspricht direkt der Biegung. Die Schallübertragung von der Membran zum Kristall kann direkt oder mittelbar erfolgen. Beachtliche Anwendungsbreite erlangte das **dynamische Mikrofon** als Umkehrung des dynamischen Lautsprechers bzw. der Kopfhörer (e). Die Membran ist dabei auf die Kalotte (auch Dom genannt) reduziert. Es gibt Ausführungen bis zum Durchmesser von 1 cm herab.

Heute ist das **Kondensatormikrofon** (1906) am wichtigsten. Einer festen Elektrode steht – in geringem Abstand isoliert befestigt – eine dünne bewegliche Membran gegenüber. Da jede eingespannte Membran Eigenschwingungen besitzt, werden in die Gegenelektrode Sacklöcher eingebracht. Ihr Durchmesser, ihre Tiefe und ihre Anzahl werden so gewählt, dass der gewünschte Frequenzgang und die Richtwirkung entstehen (f). Für einen ideal geradlinigen Frequenzgang muss das Luftvolumen auf die Steifheit der Membran abgestimmt sein. Im einfachsten Fall entsteht eine **kugelförmige Richtcharakteristik**. Es liegt ein reiner Druckempfänger vor. Durch andere Gestaltung der Löcher sind andere Richtcharakteristiken zu erreichen. Werden die Sacklöcher zu Durchgangslöchern verändert, so gelangt der Schall von beiden Seiten auf die Membran. Folglich kann nur die Druckdifferenz (d. h. die Schallschnelle) wirksam werden. So entsteht die Achtercharakteristik. Werden beide Varianten kombiniert, so entsteht die nierenförmige Charakteristik (s. u.). Durch die Verwendung von je einer isolierten Membran auf jeder Seite der Kapsel ist schließlich auch ein manuelles **Umschalten** zwischen den drei Charakteristiken möglich (g).

Den üblichen (historischen) **Betrieb** eines Kondensatormikrophons zeigt (h). Über einen sehr großen Widerstand ($R \approx 100 - 1000 \text{ M}\Omega$) wird die Kapazität des Mikrophons ($C \approx 10 - 100 \text{ pF}$) mit einer Spannung ($\approx 100 \text{ V}$) aufgeladen. Die Zeitkonstante $T = R \cdot C \approx 0,1 \text{ s}$ dann groß gegenüber den Kapazitätsänderungen durch den Schall. Daher bleibt die Ladung $Q = C \cdot U$ konstant. Folglich ändert sich die Spannung umgekehrt proportional zur Kapazität. Sie kann zum Verstärker geleitet werden. Infolge der hohen Impedanz und zur Vermeidung von störenden Leitungskapazitäten muss er unmittelbar mit der Kapsel vorhanden sein. Optimal ist dafür ein FET-Verstärker. Schließlich kann auf die Spannungsquelle dann verzichtet werden, wenn gegenüber der Membran eine Elektretfolie eingefügt wird (k). Sie besitzt ähnliche Eigenschaften wie ein Dauermagnet, hat also eine ständig vorhandene elektrische Polarisierung (fixierte Ladung), die etwa den hundert Volt entspricht. So entfällt die von außen sonst anzulegende Spannung. In Kombination mit einem FET entsteht dann ein kleines, preiswertes und dennoch hochwertiges Standardbauelement der Unterhaltungsindustrie, das Elektret-Mikrofon (j). Insbesondere bei „drahtlosen“ (schnurlosen) Kondensatormikrophonen wird eine andere Betriebsweise gewählt. Die Kapazität des Mikrophons wird mit einer Induktivität zu einem Schwingkreis zusammengeschaltet. Er wird über einen Transistor zur Eigenschwingung angeregt (i). Infolge der Kapazitätsänderung durch den Schall wird dann die Schwingung in der Frequenz moduliert. Sie kann über eine kurze Antenne – das Kabelende, was meist lose dran hängt – abgestrahlt werden. Irgendwo in der Nähe wird es empfangen und zum NF-Signal demoduliert.

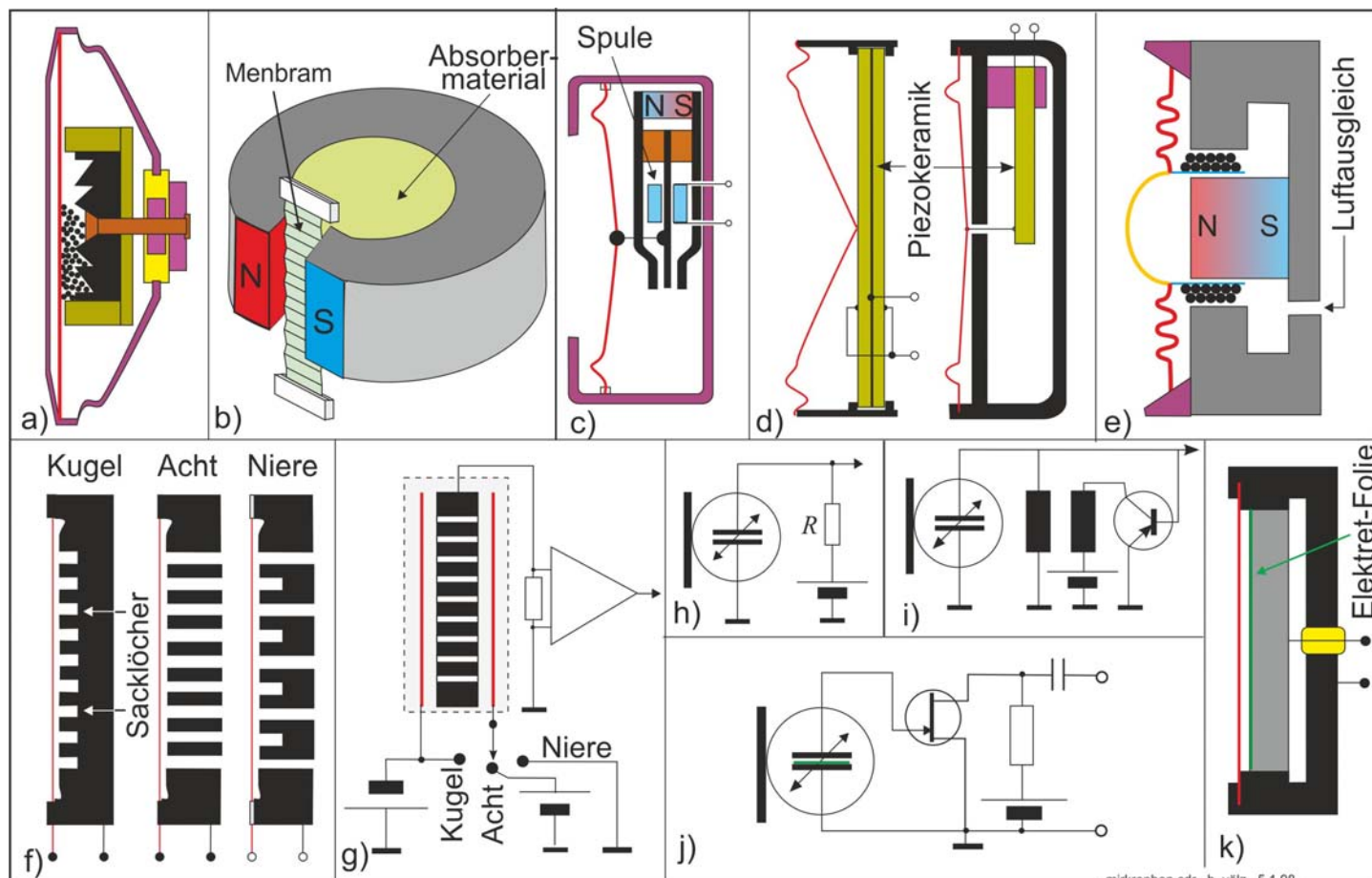
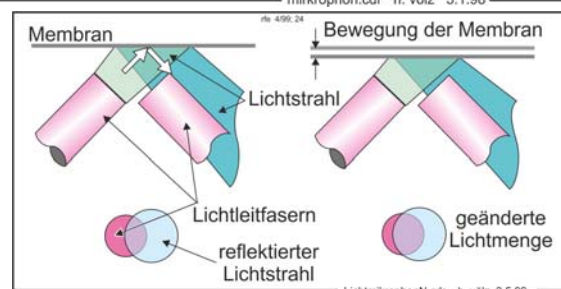


Bild 90. Die wichtigsten Mikrofonarten.

Als Besonderheit hat in den letzten Jahren Sennheiser ein **Lichtleitmikrofon** entwickelt. Durch die Bewegung einer Membran wird der Übergang zwischen zwei Lichtleitern verändert. Schließlich seien noch einige spezielle Mikrofone erwähnt: Jene für den **Körperschall**, die im direkten Kontakt mit Objekten stehen und bei Musikinstrumenten auch Tonabnehmer genannt werden. Dann die damit verwandten **Kehlkopfmikrofone** sowie die am Körper getragenen Krawattenmikrofon oder **Lavaliere** (franz. locker gebundener Künstlerknoten) erwähnt. Sehr kleine, meist unsichtbare Mikrofone heißen auch Wanzen.

Bild 91. Prinzip des Lichtleitmikrofons.



Von der **Richtcharakteristik** hängt zumindest teilweise auch der **Frequenzgang** ab. Selbst die Kugelcharakteristik hat nur bedingt einen völlig linearen Frequenzgang. Je höher die zu übertragende maximale Frequenz ist, desto kleiner muss die Kapsel gehalten werden, damit an ihr keine störenden Reflexionen auftreten. Allerdings sinkt dadurch auch die Empfindlichkeit. Für sehr hochwertige Aufnahmen sind deshalb die 1/2-Zoll-Kapseln mit rund 13 mm Durchmesser ideal.

Bei Videoaufnahmen werden häufig Camcorder mit einer **Zoomoptik** benutzt, die es ermöglicht, den Bildwinkel und damit den subjektiven Abstand von den Objekten willkürlich zu verändern. Diese Einstellung muss mit einer passenden Richtwirkung des Mikrofons kombiniert werden. Das ermöglichen besonders gut drei Mikrofone mit Nierencharakteristik (**Bild 92a**). Dabei wird das mittlere Mikrofon analog so wie ein Teil bei der Zoomoptik seitlich verschoben (vgl. Bild 64h, S. 32). So ergeben sich die Richtwirkungen von (b).

Eine besonders schmale, dafür aber weit reichende Richtwirkung ist mit einem **Parabolspiegel** (bis zu mehreren Metern Durchmesser) zu erreichen (e). Die Kombination Prinzip wird vor allem für sehr weit entfernte Objekte, z. B. bei Tieraufnahmen in der Natur oder zum „Belauschen“ bei Geheimdiensten benutzt. So wird eine sehr hohe Empfindlichkeit erreicht. Die Bündelung (Richtwirkung) des Schalls tritt aber erst dann wirksam in Erscheinung, wenn die Wellenlänge deutlich kleiner als der Spiegeldurchmesser ist.

Ferner entstand das Interferenzmikrofon, auch **Rohrrichtmikrofon** genannt (c). Vor einer Mikrofonkapsel befindet sich ein 10 bis 50 cm langes Rohr, das einseitig einen Schlitz von einigen mm Breite besitzt und geringfügig mit einem Absorptionsmaterial gefüllt ist. Dadurch ergibt sich ein Dämpfungsverlauf gemäß dem Abstand von der Membran (s. oberer Bildteil). Der einfallende Schall gelangt dann sowohl schräg durch den Schlitz als auch durch die vordere Öffnung des Rohres zur Membran. Durch die unterschiedlichen Laufzeiten (verschieden lange Wege) interferieren die Schallwellen teilweise. Dieser Einfluss nimmt mit wachsender Frequenz (kürzere Wellenlänge) zu. Insgesamt ergeben sich dadurch die Richtcharakteristiken von (d). Bei hohen Frequenzen besitzen sie neben der Hauptkeule viele Nebenzipfel, die aber im Klangbild kaum störend in Erscheinung treten.

Ganz im Gegensatz zum Rohrrichtmikrofon kann die Interferenz von Schall auch sehr störend auswirken. Das zeigt sich besonders deutlich bei schallhartem Fußboden. Dann wird an ihm die Schallquelle gespiegelt. So gelangt der Schall auf zwei auf zwei Wegen zum Mikrofon. Dabei verstärken oder löschen sich die höheren Frequenzen je nach dem Laufzeitunterschied periodisch aus. Es entsteht der stark wellige Frequenzgang. Ein Fußboden mit gut absorbierendem Material (Teppich usw.) wurde das unterdrücken. Mittelbar ist das auch zu erreichen, wenn das Mikrofon unmittelbar am Fußboden angebracht wird. Derartige Spezialmikrophone heißen **Grenzflächenmikrophone**. Für ein ausgewogenes Klangbild ist die Größe der Grenzfläche wichtig. Der zugehörige Durchmesser ($\approx \lambda/4$) bestimmt die tiefste Frequenz. Die folgende Tabelle zeigt, dass dafür erhebliche Abmessungen erforderlich

Grenzfrequenz in Hz	20	50	100	200	400
Notwendiger Durchmesser in m	5	3	1,6	0,75	0,30

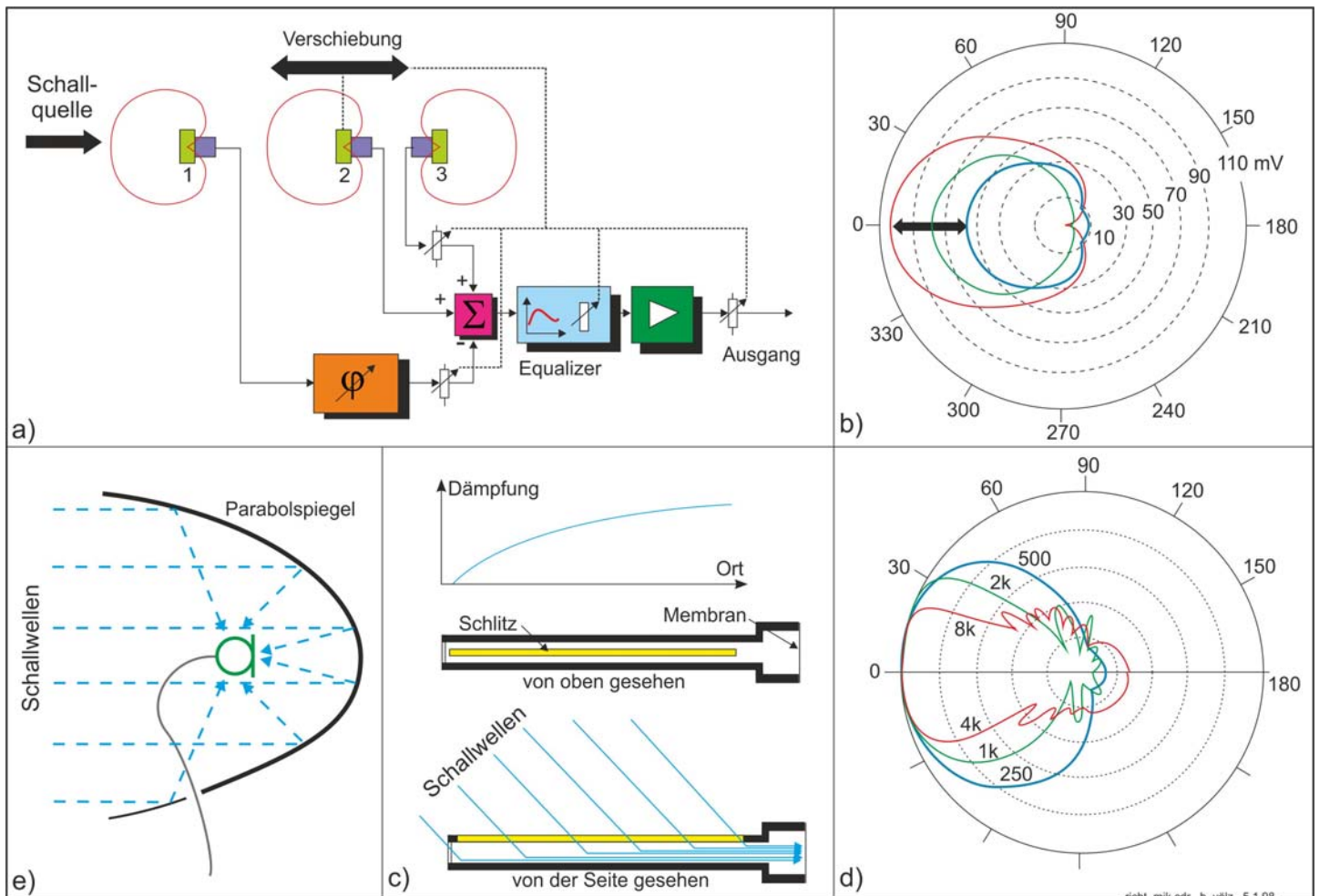
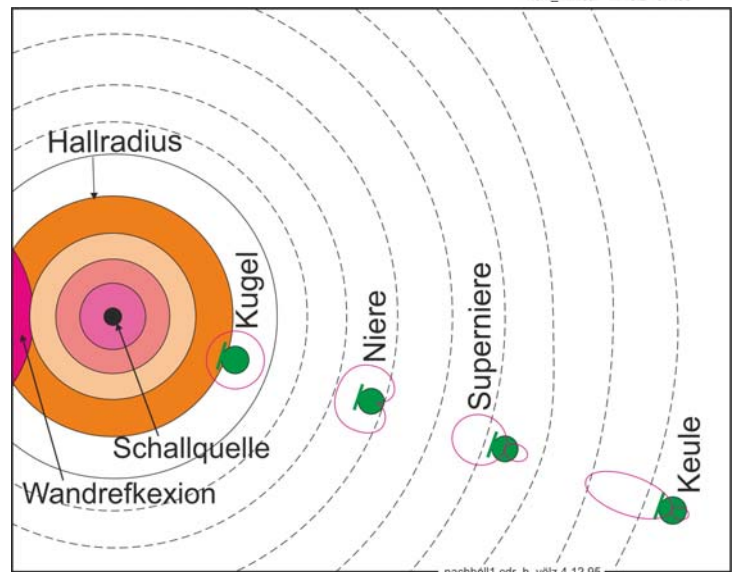


Bild 92. Ausgewählte Richtmikrofone.

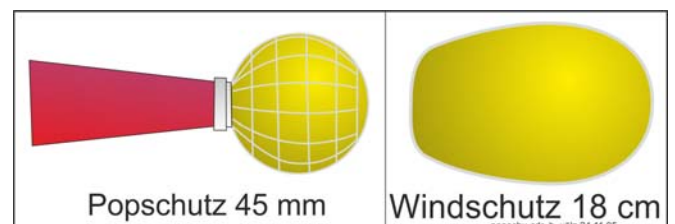
Mit jedem Schall nehmen wir auch vielfältige, typische Eigenschaften des **Raumes** wahr. Sie sind vor allem durch Reflexionen an den Wänden, der Decke und im Raum vorhandener Objekten gegeben. Von einem kurzen Schallimpuls erreicht uns zuerst der **Direktschall**, die erste Schallfront. Durch die bestimmen wir die Richtung der Schallquelle. Dann folgen einige frühe, meist einmalige **Reflexionen**, welche Beschaffenheit und Geometrie des Raumes kennzeichnen. Schließlich erreichen uns viele Mehrfachreflexionen, die ineinander übergehen und ein diffuses Klangbild vermitteln. Sie sind durch Raumvolumen bestimmt. Die **Nachhallzeit** ist dadurch bestimmt, dass diese Schallintensität auf 1/1000-tel (-60 dB) der direkten Schallintensität abgeklungen ist. Für die Aufnahmetechnik ist der Abstand von der Schallquelle wichtig. Allgemein wird angenommen, dass der direkte Schall etwa quadratisch oder kubisch mit der Entfernung von der Schallquelle abnimmt. Die Ursache hierfür ist die Absorption der Luft. Beim **Hallradius** sind Direkt- und Hall etwa gleich laut. Je nach Richtwirkung des Mikrofons ist es in Entfernungen aufzustellen, die mit dem Hallradius zusammenhängen (**Bild 93**).

Bild 93. Optimale Mikrofonabstände bezüglich des Hallradius.



Mikrofone müssen teilweise gegen störende Schallerscheinungen geschützt werden. Für den **Trittschall** hilft eine **federnde Aufhängung** des Mikrofons und **Baßabsenkung**. Der **Naheffekt**, genauer **Nahbesprechungseffekt** entsteht dadurch, dass in unmittelbarer Nähe von Mikrofonen mit Nieren- und Achtercharakteristik eine fehlerhafte Verstärkung der tiefen Frequenzen erfolgt. In etwa 5 cm Abstand entstehen so typische **Popgeräusche**. Sie werden bevorzugt durch Explosionslaute ausgelöst. Ähnliches gilt auch, besonders bei kleinen Mikrofonen, für den Wind. Hier tritt jedoch stärker ein Rauschen auf. Der **Pop-** bzw. **Windschutz** (**Bild 94**) reduzieren die Effekte deutlich. Sie bestehen aus einer etwa 2 mm starken schalldurchlässigen Schaumstoffschicht. dadurch erhält das Mikrophon indirekt eine größere Oberfläche. So werden die statischen Schwankungen des Schalldrucks über eine größere Oberfläche gemittelt werden. Daher ist der **Windschutzkorb** auch deutlich größer als der Popschutz.

Bild 94. Pop- und Windschutz.



Kopfhörer

Bei der elektrischen Übertragung muss der Schall schließlich den Menschen erreichen. Analog zu den Lautsprechern existieren hierfür mehrere Varianten. Um Störungen anderer zu vermeiden und für unterwegs sind u. a. die Kopfhörervarianten von **Bild 95** entstanden. Recht einfach ist die Methode nach (a), die dem Lautsprechern von Bild 85e, S. 40 entspricht. Die Piezokearmik kommt in (b) analog zu Bild 85i zur Anwendung. Für

hochwertige Kopfhörer wird häufig (c) gemäß dem dynamischen Lautsprecher von Bild 85f benutzt. Weiter seien noch die Elektretmembran bei Kondensatoren (d) in Bezug zu Bild 85k und die mit Spulen versehene Membran (e) in Analogie zu Bild 85l genannt.

Für den Kontakt der Kopfhörer zum Ohr werden die drei Arten von **Bild 86a** unterschieden. Der **offene** Hörer lässt zusätzlichen Umweltschall ins Ohr gelangen. Bei ihm geht also der akustische Kontakt zur Umgebung kaum verloren, wird jedoch erheblich durch den Kopfhörerschall verdeckt. Weitgehend wird er dagegen mit der **geschlossen** Variante eingeschränkt. Den Mittelweg beschreitet der **halboffene** Betrieb. Doch es gibt auch Antischallhörer, die nahezu vollständig den Kontakt zu Umwelt unterdrücken. Dazu wird zusätzlich **Antischall** erzeugt. Den offenen Betrieb wiederholt detaillierter (b). In (c) wird ein zusätzliches Mikrofon für den Außenschall eingefügt. Dieses Signal wird dann verstärkt und phasenverschoben dem externen Signal S_M so hinzugefügt, dass wird das direkt zum Ohr gelangende Signal S_D weitgehend kompensiert. Zuweilen wird das Zusatzmikrofon auch nahe dem oder im Gehörgang eingefügt (d). Dann erfolgt eine passende Gegenkopplung für das direkte Signal S_D . Formal können beide Varianten der Schallkompensation regelungstechnisch betrachtet werden. Infolge der Schallwege funktioniert das Prinzip sehr gut nur für tiefe Frequenzen und ab etwa 1 kHz geht die Wirkung weitgehend verloren ($\lambda \approx 3,4$ mm bei 1 kHz). Dadurch können immer hohe externe Warnsignale gehört werden. Deshalb ist dieser Hörer auch für Kraftfahrer zugelassen und bei den lautstarken Lastern sehr beliebt.

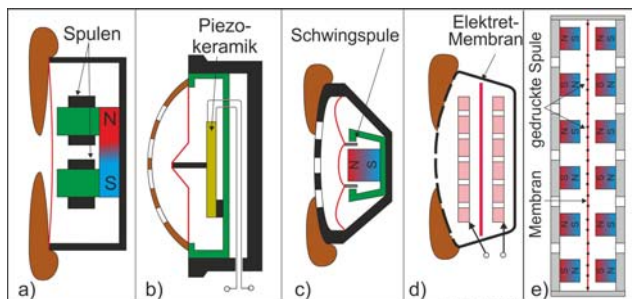


Bild 95. die typischen Kopfhörervarianten,

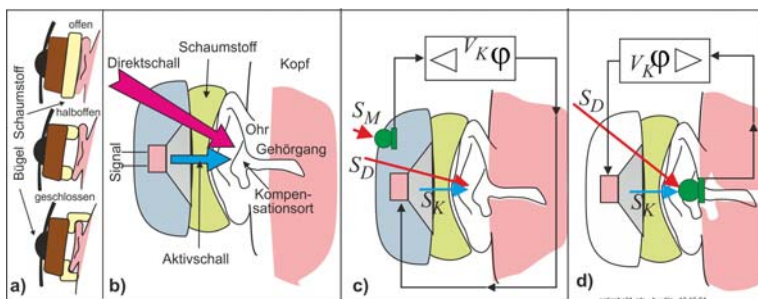


Bild 96. Betrieb der Ohrhörer (a, b) und Kopfhörer mit Antischall

Schließlich gibt es mehrere Sonderausführungen von Kopfhörern. Da sind zunächst die kleinen Ausführungen zu nennen, die mehr oder weniger tief **in den Gehörgang** gesteckt werden. Ein besonders hochwertiges System mit doppelter Bassverstärkung zeigt **Bild 97a**. Es gibt auch Surround-Kopfhörer. Doch auch bei ihnen bereitet ein echtes **Wahrnehmen des Raumes** erhebliche Probleme und das selbst bei der speziellen kopfbezogenen Aufnahmetechnik, z. B. mit Kunstkopf. Bei sehr hochwertiger Kopfhörerübertragung stört oft das Kabel. Daher entstanden die Kopfhörer mit Infrarotübertragung (drahtloser bzw. **Infrarot-Kopfhörer**). Sie enthalten dafür auch eigene Batterien. Dann gibt es die mehrere Varianten für die **virtuelle Realität**, die mit der Bildwiedergabe als Einheit getragen werden (b). Weiter müssen die extrem kleinen und meist hoch komplexen **Hörprothesen** für Schwerhörige und Hörbehinderte hervorgehoben werden. Für Sonderfälle werden schließlich auch die Knochenleitung und eine direkte **Stimulierung des Hörnervs** (Nucleus cochlearis) herangezogen.

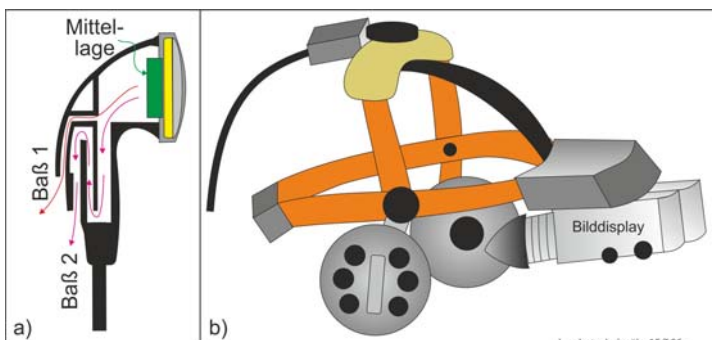


Bild 97. Innenohrhörer (a) und Headset für Schall und Bild (b)

Sonifikation

Mit Schall ist indirekt auch eine Bildwiedergabe, ähnlich dem Röntgen möglich, für die dann Sonifikation, **Sonographie**³⁰ oder photoakustische Tomografie (PAT) gebräuchlich sind. Im Gegensatz zum Röntgen treten dabei keine schädlichen Belastungen für den Patienten auf. Hinzu kommen der sehr einfache Gebrauch und die deutlich preiswerteren Geräte. Ferner entfallen aufwendige Strahlenschutzmaßnahmen und -belehrungen. Verwendeten Kontrastmittel verlassen als einzige nicht die Blutbahn. Selbst sensible Gewebe, wie bei Ungeborenen werden nicht beschädigt. Die Untersuchung verlaufen schmerzfrei. Das Prinzip und die wichtigsten Eigenschaften zeigt vereinfacht für das Herz **Bild 98**. Die erste medizinische Anwendung erfolgte 1942 durch den Neurologen Karl Dussik (1908–1968). Er stellte damit einen Seitenventrikel des Großhirns in A-Mode-Messung (Amplitudendarstellung) dar und nannte das Verfahren Hyperfonografie. Später entstanden B-Mode-Schnittbilder (bewegter Strahl, Helligkeitsdarstellung ähnlich dem Impulsradar. 1956 folgten dann Anwendungen in der Augenheilkunde und Gynäkologie. Seitdem wird die Sonifikation immer mehr und vielfältiger eingesetzt. Generell werden longitudinale elastische Schallschwingungen mit Frequenzen von etwa 2 - 20 MHz angewendet. Gemessen werden die Laufzeit (als Abstand des Reflektierenden) als Stärke des Echos und die Dämpfung der Amplitude bei Durchstrahlungen. Seit etwa 1959 sind auch Bewegungen (z. B. Blutströmungen in Adern) mittels des Dopplereffektes (ähnlich wie bei Radar) gut erfassbar. U. a. erfolgen Organuntersuchungen bei Leber, Gallenblase, Milz, Bauchspeicheldrüse, Nieren, Lymphknoten Hohlräumen (Zysten) und Steinleiden, Schilddrüse auf Vergrößerung oder Verkleinerung sowie bzgl. Tumoren. Dann krankhafte Veränderungen beim Herz und Gelenken auf Gelenkergüsse. Die Vielfalt ist so groß, dass kaum eine vollständige Erfassung möglich erscheint. Das Ungeborene kann in der Gebärmutter nahezu komplett untersucht werden, da noch keinerlei Gasüberlagerung vorliegt und die Knochenbildung erst am Anfang steht. Schlecht zu untersuchen sind u. a. Gehirn, Nerven, Wirbelsäule, Rückenmark, Inneres von Knochen und Gelenken, Luftröhre und Lunge

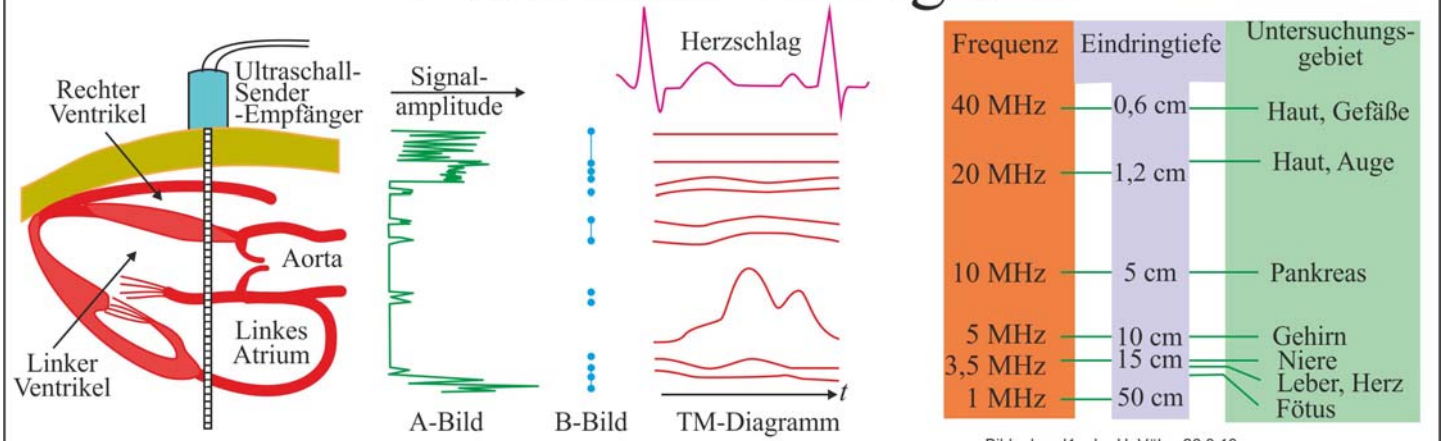
Sonderanwendungen

Der eigentliche Ursprung der Sonifikation war die militärische **Ortung von U-Boten** mit Ultraschallwellen im Wasser. Doch die damalige Energie war so groß, dass benachbarte Fische zerborsten. Erst deutlich später wurde das Verfahren zur **Aufdeckung von Materialfehlern**, wie Hohlräume, Risse usw. in Werkstoffen und Geräten eingesetzt. Auch für die **Reinigung** von Oberflächen (Optiken, Brillen usw. wird breit angewendet, ja selbst zum Zähneputzen hat sie sich bewährt. Dabei entstehen an den Oberflächen kleine Flüssigkeitsblasen, die zerplatzen und dabei den Schmutz mitreißen.

Eine deutlich andersartige Nutzung von Ultraschall erfolgte in der Anfangszeit der Rechentechnik mit **Quecksiberspeichern**. Ihr Aussehen und ihre Wirkungsweise zeigt **Bild 99**. Die zu speichernden Signale laufen verzögert in den Quecksibersäulen. So lässt sich ein ständiger Umlauf der Signale erreichen, der dann je nach Bedarf angezapft wird. Der erste dieser Speicher wurde um 1945 von Maurice V. Wilkes (*1913) und John Prosper Eckert jr. (*1919) fertig gestellt und in einem Rechner angewendet. In jeder der zehn Röhren konnten 576 Bit gespeichert werden. Später wurden bis zu 32 Röhren genutzt. Die Zeit für ein Bit betrug 1,9 μ s, die Datenrate 500 KBit/s und die mittlere Zugriffszeit 0,6 ms.

³⁰ Lateinisch **sonitus** Schall, Klang. Achtung! nicht verwechseln mit Sonogramm = grafische Schalldarstellung $\approx$ Spektrogramm.
Verstärker.doc H. Völz erste Variante 27.1.11. aktuell 29.12.2019 Seite 45 von 54

Ultraschall-Tomografie

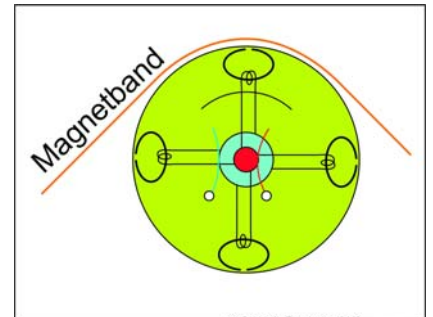


Bildgebend1.cdr H. Völz 26.3.13

Bild 98. Beispiel und Eigenschaften der Sonifikation.

Zum Hörbarmachen von Ultra- und Infraschall ist eine **Frequenztransponierung** notwendig. Heute erfolgt das mühelos digital. Doch bei der Analogtechnik war die Magnetbandspeicherung mit Änderung der Transportgeschwindigkeit erforderlich. Doch dabei wird auch die Zeitdauer der Signale verändert. Das ließ sich jedoch mit dem rotierenden Kopf nach Springer vermeiden (**Bild 100**). Durch seine Rotation wurden dann je nach Drehrichtung (Auf- oder Abwärtstransformation) zusätzlich Zeitsegmente wiederholt oder ausgelassen

Bild 100. Rotierender Magnetkopf nach Springer



rot_Springer_kopf.cdr H. Völz 19.1.08

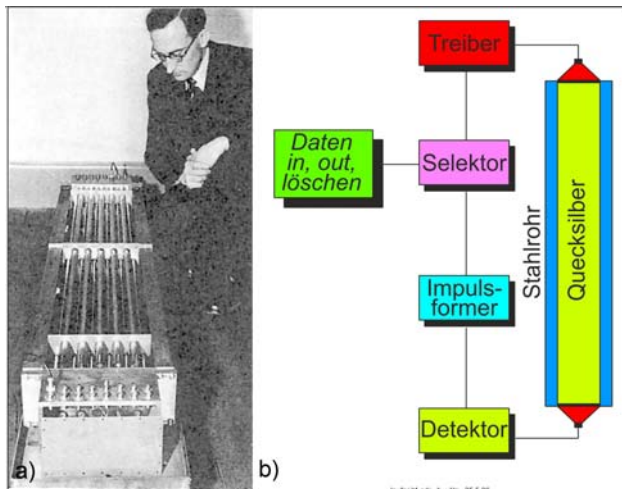


Bild 99. Quicksilver-Laufzeitspeicher mit dem Entwickler Maurice V. Wilkes (a) und seine Wirkungsweise (b).

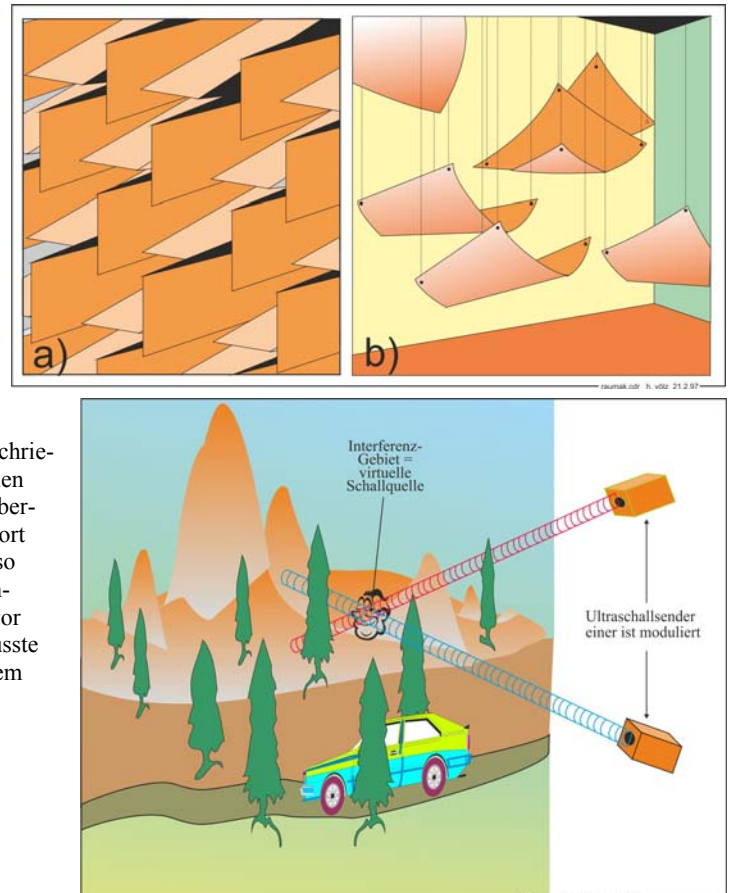
Eine sehr spezielle Anwendung ist die **Ultraschall-Holografie**. Dabei wird eine Referenzwelle von Ultraschall dem Geschehen in den Raum hinein gestrahlt. Sie interferiert mit Reflexionen an den Objekten. Die so entstehenden Schalldruckknoten werden mit einem bewegten Mikrophon gescannt und aufgezeichnet. So entsteht ein „Akustisches Hologramm“, das optisch betrachtet werden kann und sonst kaum zu gewinnende Aussagen über die einzelnen Schallquellen liefert.

Für die vielen Möglichkeiten und Aussagen zur Raumakustik, muss hier auf Spezialliteratur verwiesen werden, z. B. [Hen01], [Rei79], [Sku54] und [Vei88]. Erwähnt seien aber noch die Wirkung der Resonanzgehäuse bei Musikinstrumenten, die Schallisolierung sowie die für Messungen wichtigen „schalltoten“ und extrem lange und stark Schall reflektierenden Hallräume gemäß **Bild 101**. Das wird durch entsprechende Absorbermatten (a) bzw. Hallplatten erreicht.

Bild 101. Typisches Aussehen eines schalltoten und des Hall-Raumes.

Am Letzten sei der sonst kaum behandelte **virtuelle Lautsprecher** kurz beschrieben (**Bild 102**). Treffen zwei starke gleichartige kohärente Ultraschallquellen aufeinander, so können sie eine Nichtlinearität der Luft erzeugen. In der Überlappungszone entstehen dabei Kombinationsfrequenzen. Sie entsprechen dort einer virtuellen Schallquelle. Wird dann eine Ultraschallquelle moduliert, so entsteht dort ein virtueller Lautsprecher. Schall tritt dort quasi von einer unsichtbaren Schallquelle ausgehend auf. Das Prinzip hat Sennheiser schon vor einigen Jahren demonstriert. Entsprechend dem Reziprozitäts-Theorem müsste es auch möglich sein, analog ein unsichtbares virtuelles Mikrophon an jedem Ort einzusetzen. Hierzu ist mir leider nichts bekannt?!

Bild 102. Das Prinzip des virtuellen Lautsprechers.



virt_kompresor.cdr H. Völz 13.8.2010

7. Beginn, erster Versuch der neuen Konzeption

In etwa gilt die folgende, noch vorläufige Gegenüberstellung. Die Reihenfolge wird nur mit den ausgewählten Beispielen nur annähernd eingehalten. Dabei wird auch bereits Behandeltes nicht wiederholt. Zunächst bleiben u. a. noch unbehandelt:

Kettenreaktion, Auslöseeffekte, Katastrophen, Resonanz, Fraktale, Rekursion, usw. Einsatz von Wandlersystemen

Art, Etappe	Wo fixiert	Veränderung/Verstärkung
Ständigkeit	Realität.	Gesetze und Konstanten.
physikalisch-chemisch	Stoff (Welle und Felder), Kosmos, Erde, Objekte usw.	Hebel, Impuls, Resonanz, Strahlungen, Foto-Entwickler, Katalysator Kontrastmittel, Kettenreaktionen, Auslöseprinzip, Poincaré, Synchronisation
egotrop	Unbekannt, später Immunsystem.	Nichts Genaueres bekannt, entfällt daher hier.
genetisch	DNS-Sequenzen (Chromosomen).	Immunsystem, Antikörper, Hormon, Enzym Genetik, Bio-Rhythmen, Räuber-Beute
neuronal	Gedächtnis (Neuronen und Synapsen).	Geschmacksverstärker, Droge, Gewürz, Hypnose, Konditionierung, Lernen. Motivierung Kosmetik Parfüm + Desodorierung Reflex, Rhythmus-Takt.
gesellschaftlich	Verteilt über Individuen (Meme?).	Medienwirkungsforschung, Bestechung
technisch	Speichermaterial, Werkzeug.	Klebstoff Wellblech, -pappe, Elektronik Relais, Röhre, Transistor, Regelung, Linsen, Schall, Holografie, Fotografie, Fernrohr, Mikroskop, Tomografie, usw.
vernetzt	Internet, Expertensystem, KI	Räume und Geschehen virtuell, Fraktale, Iteration, Rekursion,

Physikalisch-chemisch

Katalysator³¹

Er ist ein Stoff, der die Reaktionsgeschwindigkeit durch die Senkung der Aktivierungsenergie einer chemischen Reaktion erhöht, ohne dabei selbst verbraucht zu werden. Ähnlich beschleunigt er die Hin- und Rückreaktion und ändert die Kinetik der chemischen Reaktionen, aber nicht deren Thermodynamik. Bei der chemischen Reaktion bildet er eine intermediären Stufe mit den Reaktanten. Nach Entstehung des Produkts wird der Katalysator unverändert freigesetzt. So kann er den Zyklus viele Male durchlaufen. Biochemische Prozesse werden durch **Enzyme** katalysiert.

Bereits in der Antike erfolgten chemische Reaktionen mit Katalysatoren. 1835 erkannte Jöns Jakob Berzelius: Viele Reaktionen erfolgen nur bei Anwesenheit eines bestimmten Stoffs, der dabei aber nicht verbraucht wird. Er nannte ihn Katalysator. Das Verständnis der thermodynamischen Ursachen gewann 1895 Wilhelm Ostwald und erhielt dafür den Nobelpreis für Chemie. Ein Katalysator verändert die Aktivierungsenergie gemäß **Bild 103**. Es wird ein „andere Weg“ – rot statt schwarz – durchlaufen. So sollte eigentlich die Verbrennung von Wasserstoff mit Sauerstoff automatisch ablaufen. Jedoch wegen der hohen Aktivierungsenergie ist die Reaktion so stark gehemmt, dass sie nur sehr geringfügig erfolgt. Ein Platin-Katalysator erniedrigt sie stark, dass sie auch bei niedrigeren Temperaturen hinreichend schnell abläuft. Bei Gleichgewichtsreaktionen verändert ein Katalysator gleichermaßen die Hin- und Rückreaktion, so dass die Lage des Gleichgewichts nicht verändert wird, sich aber das Gleichgewicht schneller einstellt.

Fast alle lebensnotwendigen chemischen Reaktionen laufen katalysiert ab (z. B. Photosynthese, bei der Atmung oder Energiegewinnung aus der Nahrung). Die meist benutzten Katalysatoren sind Eiweiße, wie Enzyme. Auch bei technischen Anwendungen sind sie oft unentbehrlich. Einige wichtige katalytische Verfahren zeigt die Tabelle:

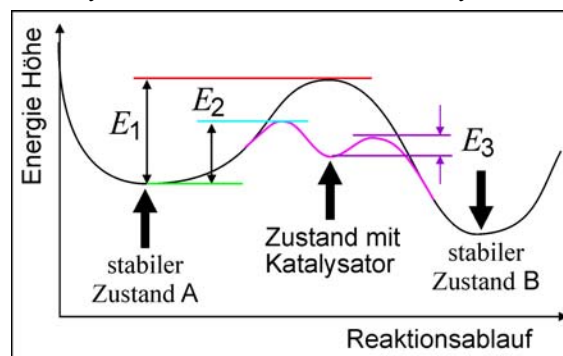


Bild 103. Zur Wirkung eines Katalysators.

Verfahren	Produkt	Katalysator	Bedingung
Haber-Bosch-Verfahren	NH_3	$\alpha\text{-Eisen}/\text{Al}_2\text{O}_3$	$T = 450 \dots 500 \text{ }^\circ\text{C}$; $p = 25 \dots 40 \text{ MPa}$
Methanolherstellung	CH_3OH	$\text{CuO}/\text{Cr}_2\text{O}_3$, $\text{ZnO}/\text{Cr}_2\text{O}_3$ oder CuO/ZnO	$T = 210 \dots 280 \text{ }^\circ\text{C}$; $p = 6 \text{ MPa}$
Kontaktverfahren	H_2SO_4	$\text{V}_2\text{O}_5/\text{Träger}$	$T = 400 \dots 500 \text{ }^\circ\text{C}$
Ostwaldverfahren	HNO_3	Platin/Rhodium	$T = 800 \text{ }^\circ\text{C}$

Fotografische Entwicklerflüssigkeit

Sie wird (auch **Entwickler** genannt) und macht die latenten Bilder eines belichteten Filmes (Fotoemulsion) sichtbar. Auf dieses Reduktionsmittel reagieren belichtete und unbelichtete Silberhalogenid-Kristalle unterschiedlich. Der Entwickler oxidiert dabei, gibt Elektronen ab, die von den Silber-Ionen aufgenommen werden und dabei zu metallischem, also sichtbarem, Silber werden. Belichtete Silberhalogenid-Kristalle bilden dabei Spuren metallischen Silbers, das dann als Katalysator die chemische Reduktion der Silber-Ionen bewirkt. Bei der anschließenden Fixierung werden die überflüssigen Reste ausgewaschen. Bei **Farbfilmen** binden sich während der Entwicklung Farbkuppler an das Silber, welches dann durch das Bleichbad aus der Emulsion entfernt wird, so dass nur noch die Farbstoffe im Film enthalten sind. Die meisten Entwickler enthalten zusätzlich Konservierungsstoffe, Säuren und Basen zur Regulierung des pH-Werts. Farbentwickler sind erheblich komplexer. In der **Holographie** werden spezifisch zusammengesetzte Entwickler eingesetzt.

Kontrastmittel

Sie verbessern die Darstellung von Strukturen und Funktionen des Körpers bei **bildgebenden** Verfahren wie Röntgen, MRT und Sonografie. Ihre Wirkung vergrößert das Signal der Untersuchung und ermöglicht teilweise zusätzliche Aussagen. Sie werden unterschiedlich verabreicht, z. B. injiziert. Teilweise erzeugen sie auch unerwünschte (schädliche) Nebenwirkungen. Daher unterliegt ihre Anwendung dem Arzneimittelgesetz. Sie müssen von **Tracern** bzw. **Radiopharmaka** unterschieden werden, die bei physiologischen Vorgängen eingesetzt werden.

³¹ Katalyse griechisch *katálysis*, *katályein* auflösen, senken, Auflösung plus lateinische Endung

Genetisch

Enzym³²

Es besteht biologischen Riesenmolekülen und wurde früher auch Ferment³³ genannt. Mit Ausnahme der katalytisch aktiven RNA (Ribozym) und DNA (Desoxyribozym) sind es Proteine. Sie beschleunigen katalytisch biochemische Reaktionen insbesondere Stoffwechselfunktionen von der Verdauung bis hin zur Transkription (z. B. Glycolyse, Citrat-Zyklus, Atmungskette, Photosynthese, Transkription, Translation und DNA-Replikation. Dabei können auch Bakterien mitwirken (z. B. Hefe). Die katalytische Wirksamkeit ist durch das aktive Zentrum bestimmt, das aus gefalteten Teilen der Polypeptidkette oder reaktiven Nicht-Eiweiß-Anteilen (Kofaktoren, prosthetische Gruppen) besteht. Entsprechend dem Schlüssel-Schloss-Modell verbinden sich die komplementären Raumstrukturen und es entsteht so der aktive der Enzym-Substrat-Komplex mit nicht-kovalente Wechselwirkungen, Wasserstoffbrücken, elektrostatische Wechselwirkung oder hydrophobe Effekten. Die Reaktionsgeschwindigkeit (Volumen pro Zeit, mol/s) hängt u. a. von Salzkonzentration, pH-Wert, Konzentrationen des Enzyms, der Substrate sowie von Aktivatoren oder Inhibitoren ab (teilweise als Sättigungshyperbel). In einem begrenzten Temperaturbereich bewirkt eine Erhöhung um 5–10 °C eine Verdoppelung. Ein Herabsetzen (Enzymhemmung bzw. Inhibition) kann durch einen Hemmstoff (Inhibitor) erfolgen. Eine Homöostase bewirkt konstante Vorgänge bei wechselnder Umwelt. Hormone wirken als externe Kontrolle dabei kann auch Reizaufnahme und -weitergabe wichtig sein. Die auch die Hormone und das Immunsystem beeinflussen.

Auch In der Medizin sind Enzyme wichtige. In der Diagnostik entdecken sie Krankheiten. So wird im Teststreifen für Diabetiker über den Blutzucker (Glucose) ein sichtbares Bild erzeugt. Enzymatisch können im Blut auch Alkohol sowie Schädigungen bestimmter Organe gefunden werden (Bauchspeicheldrüsenentzündung). Viele Arzneimittel hemmen oder verstärken die Wirkung von Enzymen zur Heilung einer Krankheit. Manche Enzymdefekte werden genetisch vererbt. Enzyme sind schließlich auch wertvolle Werkzeuge der Biotechnologie: Von der Käse-Herstellung (Labferment, aus Kälbermägen) bis hin zur Gentechnik. Gezielt wurden leistungsfähigere Enzyme (Protein-Engineering) als neuartige katalytisch aktive Proteine und katalytische Antikörper (Abzyme) entwickelt. Teilweise werden sie auch industriell u. a. bei Wasch- und Geschirrspülmitteln als Lipasen (Fett spaltend), Proteasen (Eiweiß spaltend) und Amylasen (Stärke spaltend) zur Erhöhung der Reinigungsleistung hinzugefügt. Viele Enzyme können heute mittels gentechnisch veränderter Mikroorganismen hergestellt werden. Enzyme in rohen Ananas, Kiwis und Papayas verhindern das Erstarren von Tortengelatine. Beim Obst- und Gemüseschälen werden pflanzliche Zellen verletzt und so Enzyme freigesetzt, wodurch sie vom Luft-sauerstoff braun werden. Ein Zusatz von Zitronensaft wirkt dabei als Gegenmittel. Die folgende Tabelle gibt einen Überblick zu den Einsatzgebieten der Enzyme. Ihre Wirkung wird in sieben Klassen eingeteilt: Oxidoreduktasen, Transferasen, Hydrolasen, Lyasen, Isomerasen, Ligasen (Synthetasen) und Translokasen.

technischer Prozess	Enzyme	Wirkung
Stärkeverarbeitung	α-Amylase, Glucoamylase	Stärkehydrolyse
Racematspaltung	L-Acylase	Herstellung von Aminosäuren
Waschmittel	Proteasen, Lipasen	Hydrolyse von Eiweißen, Fetten
Käseproduktion	Proteasen	Milchgerinnung
Brennerei-Produkte	α-Amylase, Glucoamylase	Stärkeverzuckerung
Brauereindustrie	α-Amylase, Glucoamylase, Proteasen	Maischprozess
Fruchtsaftverarbeitung	Pektinasen, α-Amylase	Hydrolyse der Pektine bzw. von Stärke
Backwarenherstellung	α-Amylase, Proteasen, Pentosanase	teilweise Hydrolyse von Mehl- und Teiginhaltstoffen
Lederverarbeitung	Proteasen	Weichen, Enthaaren von Leder
Textilindustrie	α-Amylase	Stärkehydrolyse, Entschlichten

Enzyme werden seit mehreren tausend Jahren benutzt. So brauten bereits die Sumerer 3000 v. Chr. Bier (alkoholische Gärung), backten Brot (Sauerteig, Hefe) und stellten Käse her. Für die Einleitung der nötigen Gärung kam im 15. Jh. die Bezeichnung „Fermentation“ ins Deutsche. Um 1890 postulierte Emil Fischer das Schloss-Schlüssel-Prinzip. 1926 zeigte James B. Sumner: das Enzym Urease ist nur ein Protein ist, das er sogar kristallisieren konnte. Für weitere Ergebnisse von 1930 erhielten 1946 John H. Northrop und Wendell M. Stanley den Nobelpreis für Chemie. Die erste Aminosäuresequenz, die von einem Enzym komplett entschlüsselt war, ist die der Ribonuklease. In den 1980er Jahren wurden katalytische Antikörper mit Enzymaktivität von Richard Lerner entdeckt.

Hormon³⁴

Es ist ein biochemischer Botenstoff, der von speziellen Zellen in Drüsen produziert wird: Zirbel-, Schild- und Bauchspeicheldrüse sowie Hypophyse. Nebennieren usw. oder Zellgeweben. Es wird überwiegend im Blutkreislauf (endokrin) weitergeleitet und wirkt so über beachtliche Entfernungen und löst an den Erfolgsorganen auch in kleinen Mengen spezifische Wirkungen aus. Chemisch sind Hormone niedermolekulare Verbindungen z. T. Peptide.

Ab dem 18. Jh. wurden in den Apotheken Substanzen aus Organen, Organsäften und Ausscheidungen von Tieren als Arznei verwendet. 1830 unterschied Johannes Müller zwischen Drüsen mit innerer und äußerer Sekretion. Später wurden Extrakte aus Nebennieren, Schilddrüsen, Ovarien und Stierhoden therapeutisch benutzt. 1901 wurde das erste Hormon (Adrenalin) entdeckt und drei Jahre später synthetisiert. 1905 wurde von Ernest Starling und William Maddock Bayliss der Begriff Hormon geprägt. Z. T. werden auch die von Nervenzellen erzeugten Neurotransmitter als Hormone bezeichnet. Zytokine regulieren das Wachstum und die Differenzierung von Zellen. Häufig treten Folgeketten von Hormonen (hormonelle Achsen bzw. Kaskaden) auf. Typische Hormonwirkungen betreffen den Stoffwechsel (u. a. Nahrungsaufnahme), Homöostase, Sexualentwicklung und -verhalten, Menstruationszyklus, Knochenwachstum, Muskelaufbau und geistige Aktivität (Anpassung an Angst und Stress). Es gibt etwa 50 verschiedene Hormone.

Die bei **Pflanzen** vorkommenden Hormone werden als Phytohormone bezeichnet. Bei **Tieren** gibt es Pheromone als Botenstoffe zwischen den Individuen. Sie sind nicht an den Organismus gebunden, in dem sie gebildet wurden und können auch über große Distanzen wirksam werden. Zuweilen werden auch **Umwelthormone** benannt. Ein Beispiel sind die Hormone der Anti-Baby-Pille, die von Kläranlagen nicht abgebaut werden. Sie werden mit dem gereinigten Wasser in die Flüsse eingeleitet. Mehr als 180 der 3000 in Deutschland zugelassenen Wirkstoffe lassen sich in deutschen Gewässern nachweisen, darunter Lipidsenker, Schmerzmittel, Antibiotika bis hin zum Röntgenkontrastmittel. Auch bestimmte Schadstoffe wie beispielsweise DDT, PCB, PBDE oder Phthalate wirken wie Hormone und beeinflussen etwa die immer früher einsetzende erste Monatsperiode bei Mädchen.

³² Es wurde 1878 von Wilhelm Friedrich Kühne als neoklassisches Kunstwort eingeführt, von altgriechisch *énzymon* abgeleitet von *en-*, *in-*, und *zýme*, was Sauerteig oder Hefe bedeutet, als die Gärung auslösender oder beeinflussender Stoff.

³³ lateinisch *fermentare* gären, machen.

³⁴ von altgriechisch *horman*, antreiben, erregen.

Immunsystem³⁵

Es ist das biologische Abwehrsystem höherer Lebewesen, kann weitgehend Schädigungen durch Krankheitserreger verhindern und entfernt dazu in den Körper eingedrungene Mikroorganismen wie Bakterien, Viren, Pilze, Parasiten (wie Bandwürmer) und fremde Substanzen. Auch fehlerhaft gewordene körpereigene Zellen kann es zerstören. Bereits einfache Organismen besitzen zum Schutz eine angeborene Variante. Selbst Pflanzen verfügen über ähnliches Verhalten. Die Wirbeltiere entwickelten zusätzlich eine komplexe, adaptive und noch mehr wirksame Immunabwehr (früher erworben genannt). Die angeborene findet innerhalb von Minuten statt und ist durch die Erbinformation lebenslang genau festgelegt. Die adaptive erfolgt zwar langsamer, ist dafür aber anpassungsfähig gegenüber neuen oder veränderten Krankheitserregern. Dabei sind besondere Zellen fähig, spezifische Strukturen (Antigene) der Angreifer zu erkennen und gezielte Abwehrmechanismen, insbesondere molekulare Antikörper zu erzeugen. Dabei wirken drei Zelltypen zusammen: die das Antigen präsentierenden (APC) die T- und die B-Lymphozyten. Nach erfolgreicher Abwehr bleiben einige Antikörper als Gedächtniszellen erhalten. Bei einer erneuten Infektion kann so deutlich schneller die Abwehrreaktion erfolgen. Das adaptive Immunsystem ersetzt aber nicht das angeborene, sondern arbeitet mit ihm zusammen. Insgesamt ist das Immunsystem ein komplexes Netzwerk aus verschiedenen Organen, Zelltypen und Molekülen;

- *Mechanische Barrieren*, wie Haut, Atemwege, Darm usw., die ein Eindringen der Schädlinge verhindern sollen
- *Lymphatisches System* aus Zellen, wie Granulozyten (in Leukozyten), natürliche Killerzellen, Lymphozyten oder Antikörper (Immunglobuline), die hauptsächlich in den Blutgefäßen und Lymphbahnen zirkulieren.
- *Proteine* als Botenstoffe oder zur Abwehr von Krankheitserregern.
- *psychische* Immunfaktoren.

Mit dem Lebensalter verringert sich die Bildung von B- und T-Lymphozyten. Dadurch erhöht sich die Anfälligkeit gegenüber Krankheiten. Ferner können auch beim Immunsystem Fehler entstehen, die nur eine schwache oder gar fehlende Immunantwort ermöglichen. Es können auch Autoimmun-Erkrankungen, (wie Multiple Sklerose Allergie (Heuschnupfen) auftreten. Die erworbene Immunschwäche AIDS wird durch das HI-Virus ausgelöst. Es können sogar Zellen des Immunsystems entarten und eine Krebserkrankung auslösen. Ferner gibt angeborene Immun-Defekte, wie Morbus Behçet und DiGeorge-Syndrom. Selbst depressive Störungen, Stress und anderen psychischen Erkrankungen können sich auf das Immunsystem auswirken. Jedoch eine Impfung kann die Effektivität und Reaktionszeit bei spezifischen Erkrankungen wesentlich zu steigern. Dagegen ist eine Unterdrückung der Immunreaktion (Immunsuppression) bei Organtransplantationen notwendig.

Reflex³⁶

Er betrifft unbewusste, organisch, neuronal festgelegte, gleichförmig ablaufende motorische Reaktion auf bestimmte äußere oder innere Reize. Die Reizantwort ist dabei festgelegt und kann angeboren oder erworben sein. Die **gelernten, erworbenen Reflexe** werden auch bedingt oder konditioniert genannt. Sie wurden erstmalig von Iwan Petrowitsch Pawlow (1849–1936) untersucht und wird auch Klassische Konditionierung (Lernen) genannt. Die **angeborenen, unbedingten** Reflexe sind bei der evolutionären Anpassung entstanden und können genetisch verankert sein. Sie können bereits von Geburt an wirksam sein oder erst zu einer bestimmten Zeit (u. a. Geschlechtsreife) wirksam werden. Beide Varianten setzen die Fähigkeit des Organismus voraus, Reize über Sinnesorgane wahrzunehmen und so zu verarbeiten, dass Muskeln betätigt werden. Beim **Eigenreflex** treten der auslösende Reiz und die Reaktion im selben Organ (meist Muskel) auf. Ein Beispiel ist der Knie- oder Patellarsehnen-Reflex. Er wird mit einem kurzen Schlag knapp unterhalb des Knies ausgelöst. Er soll spontan Gefahren reduzieren. Ähnliches können auch **Fremdreflexe** leisten. So bewirkt die Reizung der Hornhaut des Auges – etwa durch einen Luftzug – das sofortige Schließen des Augenlids. **Koordinierte Reflexe** existieren, wenn durch einen Reiz eine Gruppe von Muskeln aktiviert wird und betrifft teilweise auch Organe, wie Drüsen Herz oder Darm. Beispiele sind der Saugreflex und der Greifreflex des Säuglings

Neuronal

Die Verstärkung betrifft hierbei hauptsächlich alle Möglichkeiten, die unsere Sinne verstärkt beeinflussen.

Gewürze³⁷, den Geschmack betreffend

sind frische, getrocknete oder bearbeitete Pflanzenteile, die infolge ihres natürlichen Gehalts an Geschmacks- und Geruchsstoffen (ätherische Öle) als würzende oder den Geschmack beeinflussende Zutaten bei der Zubereitung von Speisen und Getränken aller Art benutzt werden. Allgemein sind die Funktionen von Gewürzen Geschmacksanregung, Verdauungsförderung, Verhindern von Blähungen und Förderung der Konzentration, Entspannung. Typisch sind Lorbeer- und Kaffirlimettenblätter, Blüten und Blütenteile von Safran, Gewürznelken oder Kapern, Zimt-Rinden, Zwiebeln von Ingwer, Kurkuma, Meerrettich, Wasabi und Knoblauch, Früchte oder Samen von Muskatnuss, Pfeffer, Paprika, Wacholderbeeren, Vanille, Kümmel und Anis. Bei einigen Pflanzen, wie Kümmel, Zimt und Muskatnuss können *mehrere Bestandteile* der Pflanze verwendet werden. Alle sonstigen Stoffe zur Verbesserung des Geschmacks oder der Bekömmlichkeit werden als **Würzmittel** bezeichnet. Gewürze und Würzmittel werden traditionell auch zum **Haltbarmachen** von Lebensmitteln und Getränken verwendet. Ihr Duft vertreibt oft zusätzlich Vorratsschädlinge. Sie sind bereits vor über 20 000 Jahre in Syrien und Ägypten nachgewiesen. Im Mittelalter waren sie so wichtig, wie heute etwa Öl. Auf die **Heilkraft** von Gewürzen wies schon im 12. Jh. Hildegard von Bingen hin. Trotz ihrer den Geschmack verbessernden Wirkung zählen **Zubereitungen** nicht zu den Gewürzen, u. a. Senf, Currypulver, Chutney, Sojasauce, synthetische Aromen wie Vanillin, Stoffe mineralischer Herkunft, wie Kochsalz, flüssige oder nicht allein durch Trocknung hergestellte Würzmittel (Zucker, Wein, Essig, Kokosmilch, Lakritze), wässrige, saure, ölige oder alkoholische Auszüge der Aromen von Pflanzen (Rosenwasser, Mandelöl, Nelkenöl, Vanilleöl und Knoblauchöl). **Geschmacksverstärker** sind Lebensmittelzusatzstoffe (z. B. Süßstoff). Als Einzelstoffe werden sie mit den **E-Nummern** (E 6xx) bezeichnet. Mischprodukte mit einem hohen Anteil an Aminosäuren wie etwa Hefeextrakt, Hydrolysate von Proteinen oder Aromen zählen nicht dazu. Prinzipiell besteht für sie die Kennzeichnungspflicht. In Wiki sind rund 40 Stoffe drunter 20 echte aufgezählt. Als Geschmackswahrnehmungen existieren nur **salzig, sauer, süß** und **bitter**. Sie werden über die Geschmacksknospen wahrgenommen. Als 5. Geschmack tritt **umami** mit anderer Wirkungsweise auf. Er beruht weitgehend auf Glutaminsäure, dessen Verbindungen (E 620 bis E 625) bei empfindlichen Menschen Taubheitsgefühl in Nacken, Rücken und Armen, Kopfschmerzen, Herzklopfen, Schwächegefühl auszulösen können.

Parfum³⁸, den Geruch betreffend

Es geht auf die frühe Anwendung von Räucherstoffen zurück und ist ein meist flüssiges Gemisch aus Alkohol und **Riechstoffen**, das der Erzeugung angenehmer Gerüche dienen soll. Anwendungen sind:

³⁵ lateinisch immunis unberührt, frei, rein

³⁶ lateinisch reflexus das Zurückbeugen, vgl. Reflektieren von Licht.

³⁷ mittelhochdeutschen wurz (althochdeutsch wurz und wurza, mhd. wurz(e), nhd. wurz). es bedeutete ursprünglich einfach Wurzel und ist noch den Wörtern Haselwurz und Nieswurz erhalten.

³⁸ französisch parfum, lateinisch per fumum ‚durch Rauch‘.

1. **Riechstoffe** im engeren Sinn und deren Mischungen (Duftkompositionen mit bestimmtem Odeur), die den Körpergeruch verändern oder einen Odor überdecken. Sie dienen dem Wohlbefinden und der Darstellung
2. **Raumdüfte** für Innenräume, können durch Sprays oder über aufgestellte Träger im Raum eingebracht werden.
3. **Geruchsstoffe**, die in Parfüms eingebracht werden. Sie sollen den Konsumenten Produkte attraktiv machen. Das betrifft viele Produkte für Bad, Küche, Haus und Garten. Produkte mit unangenehmen starken Eigengeruch – wie Reinigungsmittel oder Haarfärbemittel – werden mit ihnen angenehm gemacht (*funktionalen Parfümerie*). In der Lebensmittelindustrie gelten *Aromastoffe* ebenfalls als unverzichtbar, z. B. Vanille, die in Süßspeisen und Parfüms verwendet wird.

Bereits in den alten Hochkulturen **Ägyptens**, besonders mit der Pharaonin Hatschepsut (1490–1469 v. Chr.) wurden Duftmischungen von speziellen Priestern aus Harzen, Balsamen und Salben hergestellt. Der Wohlgeruch wurde ein wichtiger Ausdruck des Lebens und fester Bestandteil von reinigenden Ritualen. Das wichtige duftende Kosmetik Kyphi enthielt gemischt mit Ölen, Wein und Rosinen besonders Weihrauch, Styrax amber, Zimtrinde, Opoponax, Myrrhe, Kalmus, Galgant, Benzoeharz, Oud, Sandelholz und Rosenblätter. Teilweise mussten dafür die Rohstoffe von weither geholt werden. Später wurde diese Entwicklung auch von den Arabern übernommen und den Römern genutzt. **Indien** war das wichtigste Land der Rohstoffe, die vom Himalaja im Norden bis zum Indischen Ozean wachsen. Dort wurden sie für **Räucherrituale**, parfümierte Salben und Öle, für medizinische Zwecke sowie zur Reinigung des Körpers verwendet. Mit dem Kamasutra ist sowohl die Kunst eines erfüllten Liebeslebens überliefert, als auch der Umgang mit aromatischen Substanzen. So gab es duftende Cremes für den Körper, parfümiertes Wachs für die Lippen und gründlich geputzte Zähne. Über die **Kreuzzüge** kamen Rohstoffen und Mixturen des Orients in die **abendländische Kultur**. Vorher war nur Lavendelwasser bekannt. Als im 15. Jh. Destillate von hoher Konzentration hergestellt werden konnten, kamen die ersten ätherischen Öle in den Handel. Durch Katharina von Medici (1519–1589) am Hofe Heinrich II. kommt 1580 der Alchimist und Apotheker Francesco Tombarelli nach Grasse (Frankreich) und eröffnet ein **Laboratorium** zur Herstellung von Düften. So entstand das Gründerzentrum der europäischen Parfümindustrie. Duftwässer wurden zum unverzichtbaren Hilfsmittel der täglichen Toilette (Eau de Toilette). Synthetisch hergestellte Duftstoffe entstanden mittels chemischer Forschung. Natürliche Riechstoffe werden aus zerkleinertem Rohmaterial durch Destillation, Mazeration, Enfleurage, Extraktion oder Auspressen (Expression) gewonnen.

Die **Verdünnungsklassen** sind: *Eau de Solide* (1–3 %), *Eau de Cologne*, Kölnisch Wasser (3–5 %), *Eau de Toilette* (6–9 %), *Eau de Parfum* (10–14 %), *Extrait Parfum* 15–30 % und *Intense-Variante* bis 40 %. Bei der **Duftintensität** (-wirkung) werden Schwellen unterschieden: *Duftwirkung* als noch nicht wahrnehmbare Intensität, *Wahrnehmung*: Man riecht etwas, kann es jedoch noch nicht zuordnen, *Erkennbarkeit*: Der Duft ist erkennbar und benennbar. Es folgen *angenehmer*, *aufdringlicher* Duft bis hin zur *Flucht-Schwelle*. Die Duftnoten sind: *Kopfnote* unmittelbar nach dem Auftragen; *Herznote* in den Stunden, nach dem Verflüchtigen der Kopfnote und die *Basisnote* als lang anhaltender, schwerer Bestandteil. Als **Duftfamilien** gelten u. a. französisch, zyprisch, holzig, orientalisch, ledern, gourmandisch und tropisch. Seit 1997 werden in der EU 26 Duftstoffe als potentiell allergieauslösend eingestuft. Das absichtliche oder gezielte Verdecken unerwünschter Gerüche wird **Desodorieren** genannt. Bei **kosmetischen** Produkten gegen Körpergeruch werden Desodorantien und Antitranspirantien (vorübergehende Verengung der Schweißdrüsen) unterschieden. Sie werden als Sprays (Aerosol), Pump-Sprays, Stifte, Roller und auch als Emulsionen (Crème) angeboten.

Drogen³⁹, die Psychologie betreffend

Allgemein werden damit psychotrope Stoffe und deren Zubereitungen bezeichnet, die rauschartige Zustände bewirken (Rauschdrogen, -mittel oder -gifte). Sie ändern körperliche Zustände, das Bewusstsein und/oder die Wahrnehmung. Wissenschaftlich werden jedoch nur getrocknete Teile von Pflanzen, Pilzen, Tieren oder Mikroorganismen eingeordnet, die zur Arzneimittelherstellung verwendet werden (Arzneidrogen). Einige psychoaktive Drogen werden traditionell als Genussmittel angesehen und konsumiert, z. B. Koffein (meist als Kaffee oder Tee), dann Alkohol (Bier, Wein, Schnaps), Nikotin (Tabak), Cannabis (Marihuana, Haschisch), Kokain (Kokablätter), Betel sowie Kath. Bei höherer Dosierung können sie auch das Bewusstsein ändern oder schädlichen Folgen bis hin zu Abhängigkeit und Tod verursachen. Rechtliche Regelungen zur Benutzung wurden und werden zeitlich und örtlich unterschiedlich gehandhabt. In Sonderfällen werden sie auch medizinisch eingesetzt.

Der Gebrauch von Rauschgiften ist bereits um 6 000 v. Chr. durch Weinbau in Zentralasien und ab 3000 v. Chr. durch Bier in Ägypten belegt. Auch die Anwendung von Fliegenpilzen usw. ist sehr früh belegt, Schlafmohn etwa ab 1000 v. Chr. Ritueller Gebrauch ist umfangreich belegt. Besonders häufig erfolgt der erste Kontakt mit Drogen bei Jugendlichen. Nur wenige der Erstkonsumenten gehen zu einem regelmäßigen Konsum über.

Hypnose⁴⁰, die Bewusstheit betreffend

Sie betrifft physiologisch und psychologisch einen künstlich erzeugten partiellen Schlaf in Verbindung mit einem veränderten Bewusstseinszustand (hypnotische Trance oder Induktion). Als tief entspannter Wachzustand mit extrem eingeschränkter Aufmerksamkeit, können nur wenige, vor allem geringe Aufmerksamkeit fordernden Inhalte wahrgenommen werden. So wird die Kritik gezielt umgangen und schrittweise ausgeschaltet und das Bewusstsein verliert seine beherrschende Stellung. Es werden die Selbst- und Fremdhypnose unterschieden. Ein die Hypnose Bewirkender heißt Hypnotiseur. Die medizinische Anwendung wird auch Hypnosetherapie oder -dation genannt. In den 1930er Jahren traten u. a. Ferenc M. Sándor und Erik Jan Hanussen als Hypnotiseure sehr erfolgreich auf. Vorübergehend benutzte auch Freud die Methode, ging jedoch später zum freien Assoziieren über.

Welche Suggestionen am besten geeignet sind, ist vom dem Probanden und den Umständen abhängig. Bei der direkten Variante werden vorwiegend befehlsähnlichen Suggestionen benutzt, bei der indirekten dagegen besondere Sprachmuster, die eher einen erlaubenden oder gewährenden Charakter besitzen. Es gibt auch Blitzvarianten, welche die Trance innerhalb weniger Sekunden induzieren. Sie setzen hohe Erwartungshaltung und ein Überraschungsmoment voraus und werden vor allem zur öffentlichen Schau genutzt, haben daher nichts mit der therapeutischen Hypnose zu tun.

Die Hypnotherapie ist per Kernspinresonanztomographie (MRT) und der Elektroenzephalographie (EEG) belegt und wird gegen Depressionen, Suchtkrankheiten, Sprechstörungen, zur Steigerung des Selbstwertgefühls, zum Stress- und Abbau, bei Schlafstörungen und zur Schmerzminderung (Geburt) eingesetzt. Meist ist es jedoch besser, die Ergebnisse durch den eigenen Willen zu erreichen. Ob der Proband in einer Hypnose auch zu kriminellen Handlungen bewegt werden kann, ist nicht endgültig geklärt.

Teilweise verlangt die Hypnose bereits mehrere Menschen und leitet damit zu den gesellschaftlichen Verstärkungen über.

Gesellschaftlich, mehrere Teilnehmer betreffend

Lernen⁴¹

Für alle (höheren) Lebewesen ist Lernen eine wesentliche Voraussetzung für ihre Anpassung an die (sich ändernde) Welt. Dabei erfolgen Speicherungen im neuronalen Gedächtnis (s. #####). So werden neue Einsichten gewonnen und Schlussfolgerungen gezogen. Oft erfolgt dabei eine Bewertung der Umwelt, die zur Verknüpfung mit Bekanntem (Erfahrung, Strategien) und zum Erkennen von Regelmäßigkeiten (Mustererkennung) wesentlich

³⁹ französisch *drogue* und niederländischen *droog* = trocken, niederdeutsch auch *dröge*, gelangte gegen 1600 ins Deutsche.

⁴⁰ altgriechisch *hýpnos* Schlaf.

⁴¹ Etymologisch ist es mit „lehren“, „List“ und „leisten“ verwandt. was ursprünglich „einer Spur nachgehen, nachspüren, schnüffeln“ bedeutet. Im Gotischen heißt *lais* „ich weiß“, bzw. „ich habe nachgespürt“. Lernen hinterlässt Spuren!

ist. Es ist so die Grundlage aller Erinnerungen (s. **Zeit**). Das erfolgt weitgehend nach dem Ursache-Wirkungs-Prinzip und ist daher immer ein zeitlicher, oft auch mehrstufiger Prozess. Daher ist „Lernprozess“ ein überflüssiger Pleonasmus. In diesem Sinne ist Lernen ein wichtiger Aspekt der **Informationstheorie** (s. **###**). Beim Menschen beeinflusst es sein Denken, Fühlen und Handeln und ermöglicht dabei den gespeicherten Erwerb von Wissen (s. **####**) und von Fertigkeiten auf geistigem (sprachlichem), körperlichem, charakterlichem und sozialem Gebiet.

Generell ist zwischen einem **Dazu- und Umlernen** zu unterscheiden. Wer lernt (Lernkurve), kann (und muss teilweise) auch vergessen. So kann eine Fertigkeit oder Kompetenz zum Teil oder vollständig verloren gehen. Die **Vergessenskurve** war und ist eine wichtige Grundlage zur Bestimmung der **Gedächtnisleistungen**. Beide können – je nach der Kategorie des Inhalts – verschieden sein (s. **####**). Grob werden weiter **beiläufiges** (inzidentelles, implizites), **absichtliches** (intentionales), aber auch **entdeckendes, selbstbestimmtes, expansives** und **widerständiges** Lernen unterschieden. Außerdem gibt es **individuelles und kollektives** Lernen. Beides kann mittels (verschiedener) Lernmethoden und -strategien planvoll, systematisch gestaltet werden. Lernen erfolgt auf der Grundlage **aktiver Neuronen** in Teilen des Gehirns. Die genaue Funktionsweise ist nur teilweise wissenschaftlich gesichert und sogar umstritten. So entstanden schon in der Antike erste Vorstellungen (s. **Mnemotechnik**). Insbesondere seit dem Wissen um die Neuronen wurden mehrere unterschiedliche **Lernmodelle** entwickelt. Mit der Informatik entstand auch **programmiertes Lernen**, teils mit Computern statt Lehrern. Die beim Lernen einbezogenen Hirnreale werden seit geraumer Zeit im **Papez-Neuronenkreis** zusammengefasst. Recht sicher ist, dass schon das evolutionär alte paleo- und archikortikale limbische System essentielle Lernvorgänge ermöglichte und heute noch Grundlagen höherer Gedächtnisleistungen bereitstellt. Einen Durchbruch bei einzelnen Vorgängen erreichte Eric Kandel, der dafür den Medizin-nobelpreis 2000 erhielt.

Einfache Formen des Lernens sind Habituation, Sensitivierung und Konditionierung. Sie gibt es bereits bei wirbellosen Tieren, wie der Meeresschnecke Aplysia und der Fruchtfliege Drosophila melanogaster. **Habituation** betrifft nur die Aktivierung einer Synapse. Die durch sie ausgelöste Reaktion wird bei der Wiederholung schrittweise schwächer. Es gibt Kurz- und Langzeit-Habituation. Bei beiden erfolgt eine Inaktivierung der Calcium-Kanäle in der Präsynapse. **Sensitivierung** betrifft mindestens zwei Synapsen und führt zu einer erhöhten Botenstoffausschüttung, wodurch z. B. die Beweglichkeit der Vesikel der Botenstoffe erhöht oder die Kaliumkanäle blockiert werden. Die **klassische Konditionierung** (Iwan Petrowitsch Pawlow: bedingter Reflex) ist einfachste Form des **assoziativen Lernens**. Sie erfolgt durch die Verknüpfung von Reizen, die hintereinander auftreten. Der erste löst noch keine Reaktion aus. Sie tritt erst beim zweiten ein. Werden beide Reize mehrfach wiederholt, so genügt schließlich der erste Reiz allein für die Reaktion. Es hat sich eine **Assoziation**, Verknüpfung, zwischen dem ersten bedingten Reiz (conditioned stimulus) und dem zweiten unbedingten Reiz (unconditioned stimulus) gebildet. Die **Operante Konditionierung** beeinflusst spontanes Verhalten und betrifft die positive Auswirkung von Belohnung bzw. beim Ausbleiben von Bestrafung. Sie ist bereits der Übergang zu komplexem Lernen.

Das **Spiel** ist die ursprüngliche und hoch kreative Form des Lernens bei allen höher entwickelten Tieren und beim Menschen. Spielen ist zwar nicht zweckorientiert, ist aber (gerade deshalb) für die Ausbildung und Fortentwicklung der höheren kognitiven Fähigkeiten wichtig. Bei vielen Tieren existiert die Spielphase nur zu Anfang. Tiere mit einer beachtlichen Intelligenz (wie einige Rabenvögel, Papageien, Delphine, Affen) und der Mensch spielen dagegen bis ins hohe Alter. Kinder sollten von Beginn an spielen, um so die Welt zu entdecken und sich zu eigen machen.

Auswendiglernen verzichtet auf Kenntnisse der inneren Komplexität und Schlussfolgerungen. Es ist auf die originalgetreue Wiedergabe der Lerninhalte ausgerichtet. Grundlage dafür ist die häufige (ca. 30-malige, s. **####**) Wiederholung. Eine nicht zu unterschätzende Rolle beim Lernen kommt dem **Humor** zu. **Dialogisches Lernen** betrifft einen gleichberechtigten Dialog (auf Augenhöhe) zwischen Partnern, Machtansprüchen sollten hierbei nicht auftreten. Das Sokratische Gespräch ist dabei als Ursprungsform anzusehen. Der Begriff **E-Learning** betrifft das Lernen mit Hilfe von Computer und deren Programme. Der Spezialfall **M-Learning** (von mobil) könnte beim Smartphones Bedeutung gewinnen. In beiden Fällen können erweiterte digitale Inhalte mit Kombinationen aus Text, Bild, Video und Ton genutzt werden. Bei **Enkulturation** werden kulturelle Normen, Werte und Verhaltensweisen erlernt, die für die eigene Kultur wünschenswert oder erforderlich sind. Einflüsse hierzu stammen von Eltern, anderen Erwachsenen und Gleichaltrige. **Episodisches Lernen** betrifft Verhaltensänderungen als Folge eines erlebten Ereignisses. Verschiedene Lerninhalte werden in speziellen Gedächtnissen gespeichert (s. **Gedächtnis**).

Eine **Motivierung**⁴² betrifft die Ursache, den Zweck und Beweggrund sowie Leitgedanken für den daraus folgenden Ablauf, das Ziel und die Intensität von menschlichen Verhalten. Wichtig ist sie beim Lernen und kann dann durch **Belohnung oder Bestrafung** beeinflusst werden. Unerwünscht und strafgesetzlich verboten ist die entsprechende Bestechung eines Amtsträger (Beamter, Angestellter im öffentlichen Dienst usw.). Sie gehört zur **Korruption**. Dabei wird die materielle von der immateriellen Vorteilsannahme unterschieden. In der Erzähltheorie wird zwischen einer kausalen und finalen Motivierung unterschieden.

Die **Medienwirkungsforschung** untersucht deutlich allgemeiner die Auswirkungen von Medien auf die Rezipienten und zwar sowohl auf einzelne Personen als auch auf Gruppen und Gesellschaften. Sie begann mit der Faszination der neuen Medien **Kino und Radio**. Und betraf zunächst vorwiegend die Werbung und politische Propaganda. Besonders deutlich wurde sie aber durch die Hörspiel-Adaption „Der Krieg der Welten“ von Orson Welles 1938 über eine außerirdische Invasion. Lazarsfeld stellt einen Verstärkereffekt fest: Massenmedien verändern bestehende Einstellungen kaum, sondern verstärken sie. Das gilt vor allem für die öffentliche Meinung und die Weltbilder der Rezipienten. (s. hierzu oben auch 6. bei Funk, Tomografie und Sonifikation)

Technisch

Viele Beispiele hierzu sind bereits in 2. Austausch von Parametern, 3. mit Impulsen, 3. Relais, 4. Elektrische Signale, 5. Effekte und 6. bei Schall und Bildern vorhanden. Hier werden als Beispiel nur **Klebstoff, Wellblech, Hohlrohre und Wellpappe** ergänzt

Klebstoff⁴³ (z. B. Leim, Kleister oder Kitt, früher auch Kleber) ist ein nichtmetallischer Werkstoff (z. B. Polymer), um Etwas zusammenzufügen, fest miteinander zu verbinden. Das erfolgt durch Oberflächenhaftung (Adhäsion) und innere Festigkeit (Kohäsion) ohne dabei das Gefüge der Teile wesentlich zu verändern. Zusätzlich können Klebstoffe weitere Aufgaben übernehmen, z. B. Schwingungsdämpfung, Abdichten gegen Flüssigkeiten und Gase, Ausgleich unterschiedlicher Füge-teil-Dynamiken, Korrosionsschutz, thermische und elektrische Isolation oder Leitfähigkeit. Es gibt heute eine sehr große Anzahl von Varianten. Viele Produkte, wie z. B. Bücher, Mobiltelefone, Fußböden, Matratzen, Autos usw. wären in ihrer heutigen Form ohne Klebstoffe nicht realisierbar. Andererseits ist Kleben ist eine der ältesten Kulturtechniken der Menschheit. Es ermöglichte die Herstellung von Waffen und Werkzeugen und half sich gegen eine feindliche Umwelt durchzusetzen. Einige zeigen Anwendungen vor über hunderttausend Jahre. Die erste handwerkliche Leimfabrik wurde 1690 in Holland gegründet. Heute besteht eine kaum übersehbare Vielfalt an Klebstoffen. Hier genügt ein sehr gekürzter Überblick:

Wasserlösliche Klebstoffe werden meist als Leim bezeichnet, dazu gehören die natürlichen Kasein-, Glutinleime, Stärkeprodukte (Kleister) und Naturharze. Heute sind **synthetische** Klebstoffe besonders verbreitet. **Schmelzklebstoffe** werden heiß als Schmelze aufgetragen und bilden dann beim Erstarren die Klebverbindung. Sie bestehen meist aus Ethylen-Vinylacetat-Copolymeren mit Harzen und Wachsen. Hierzu gehören auch die Heißsiegelklebstoffe, die z. B. zum Verschließen von Joghurtbechern dienen. **Kontaktklebstoffe** (vor allem zum Verkleben von Leder und Textilien) werden auf die Oberflächen aufgetragen. Nach einer Vortrocknung erfolgt die Klebung durch kurzzeitiges, starkes Zusammendrücken. Hierbei werden u. a. Copolymere von Butadien mit Styrol oder Acrylnitril benutzt. **Haftklebstoffe** werden für Haftetiketten und Dekofolien verwendet und sind meist dauerklebend auf Kautschukbasis: sie sind nur wenig haltbar und neigen bei andauernder Last zum »Kriechen«. **Alleskleber** sind vielseitig verwendbar und bestehen aus u. a. aus Lösungen von Polyurethan oder Polyvinylacetat in Estern, Ketonen oder Kohlenwasserstoffen. Sie binden durch

⁴² lateinisch movere, motum bewegen.

⁴³ althoch deutsch kleben festsitzen, festhaften.

Verdunsten des Lösungsmittels. **Anlösende** Klebstoffe werden für Kunststoffe verwendet und enthalten ein Lösungsmittel (z. B. Aceton für Polystyrol), das die zu verklebenden Oberflächen anöst. **Reaktionsklebstoffe** (Ein- und Zweikomponentenkleber) wirken durch chemische Reaktionen und enthalten u. a. Aminoplaste oder Epoxydharze. **Härtungsklebstoffe** verwenden UV-Bestrahlung zum Verkleben.

Bei den Mechanismen der Klebung werden unterschieden: **Härtung** durch Polymerisation oder Polyaddition, **Verfestigung** durch Trocknen, Abkühlen, Polykondensation oder durch Gel-Bildung, **Kombination** verschiedener chemischer Mechanismen auch in Kombination mit physikalischen.

Klebebänder, ein- und doppelseitig, **Haft-Etiketten**, beschichteter Teile usw. sind eigentlich keine Klebstoffe. Die typischen funktionalen Träger sind Papier, Polymerschäume und -folien, Metallfolien, Metall- und Kunststoffhalbzeuge.

Das **Verb kleben** bedeutet auch: An etwas unbedingt festhalten wollen (Posten, Amt), mit etwas verbunden sein, am Alten festhalten, symbolisch an seinen Händen klebt Blut, er hat gemordet, umgangssprachlich jemand eine Ohrfeige geben.

Als **Wellblech** bezeichnet man eine Platte aus Blech, die durch ihre wellenförmig Verbiegungen bei geringer Dicke des Blechs eine vergleichsweise sehr hohe Steifigkeit und Tragfähigkeit erreicht. Es besteht vorwiegend aus feuerverzinktem Stahl. Seltener ist die Beschichtung mit einer Zink-Aluminium-Legierung (Galvalume). Die ersten hergestellten Wellbleche besaßen eine nur geringe Wellenhöhe. Heute werden Trapezbleche mit Steghöhen von 20 cm Tiefe bei nur 0,5 mm Blechdicke hergestellt. Ihre Herstellung erfolgt hauptsächlich durch Walzprofilieren bzw. Breitbandprofilieren. Als Erfinder gilt Henry Robinson Palmer. Er ließ es sich 1829 patentieren. Nach dem Auslaufen des Patents 1843 wurde es zu einem, weltweit wichtigem Baumaterial. Seit etwa 1850 wurde es in England zur Dachdeckung verwendet.

Analog dazu sind **Hohlrohre** und für Verpackungen usw. **Wellpappe**. Beide reduzieren ebenso den Materialaufwand.

Vernetzt

Vor allem **Mathematik, Fraktale, Rekursion, Kettenreaktion, Computertechnik und virtuelle Welten**, muss noch erarbeitet werden

Literatur

- [Hen01] Henn, H., Sinambari, G. R., Fallen, M.: Ingenieur-Akustik. 3. Auflage. Vieweg, Braunschweig - Wiesbaden, 2001
- [Rei79] Reichardt, W.: Gute Akustik - aber wie. Verlag Technik Berlin, 1979
- [Sku54] Skudrzyk, E.: Die Grundlagen der Akustik. Springer - Verlag Wien 1954
- [Völ05] Völz, H.: Handbuch der Speicherung von Information Bd. 2 Technik und Geschichte vorelektronischer Medien. Shaker Verlag Aachen 2005
- [Völ56] Völz, H.: RC-Generatoren. Elektron. Rundschau 10 (1956) 1, 21-24; 2, 49 - 52
 dito Der RC-Resonanzverstärker. Elektron. Rundschau 10 (1956) 3, 69 - 70
 dito Der Phasenschiebegerator. Elektron. Rundschau 10 (1956) 10, 275 - 278 und 11, 308 - 309
- [Völ59] Völz, H.: Beitrag zum Verstärker mit extrem kleinem Innenwiderstand. Frequenz 13 (1959) 7, 212 - 222
- [Völ66] Völz, H.: Elektronische Spannungsstabilisation. VEB Verlag Technik Berlin 1966, 2. Aufl. 1969
- [Völ74] Völz, H.: Elektronik für Naturwissenschaftler. Akademie Verlag, Berlin 1974
- [Völ89] Völz, H.: Elektronik. 5. Aufl. Akademie Verlag, Berlin 1989
- [Völ96] Völz, H.: Kleines Lexikon - Audio - und Video - Technik, Verlag Technik, Berlin 1996
- [Völ99] Völz, H.: Das Mensch - Technik - System. Expert - Verlag, Renningen - Malsheim 1999 + Linde - Verlag
- [Vei88] Veit, I.: Technische Akustik. Vogel Buchverlag, Würzburg 1988

Inhalt

1. Einführung	2
Mögliche Arten zur Verstärkung	2
2. G-Verstärker = Austausch von Parametern	3
Hebel, schiefe Ebene, Keil und Schraube	3
Rollen und Hydraulik	4
Reib- und Zahnradgetriebe	4
Transformatoren	5
3. I-Verstärker mittels Impulsen	5
3. D-Verstärker beginnen mit dem Relais	6
4. S-Verstärker für elektrische Signale	8
Elektronenröhren	8
Die Penthode	8
Weitere Verstärkerröhren	10
Geschichte der Röhren	10
Transistoren	10
pn-Übergang	11
Transistoren	12
Bipolar-Transistor	13
Sperrschicht-Feldeffekttransistor	14
MOS-FET	15
Transistor Geschichte	16
Betriebsarten der elektronischen Verstärker	16
Rückkopplungen	18
Operationsverstärker	20
RC-Verstärker und Generatoren	21
Stabilisierung	22
Tunneldiode und Gunn-Effekt	24
5. E-Verstärker durch Effekte	24
Maser und Laser	24
Aktive Bauelemente für Mikrowellen	26
6. O-Verstärker für Bilder und Schall	29
Vielfalt optischer Linsen	30
Fotografie	31
Fernrohre und Teleskope	32
Mikroskope	34
Tomografie	36
Funk-Antennen	37
Möglichkeiten bei Schall	39
Der elektrodynamische Lautsprecher	40
Schallabstrahlung und Frequenzgang	41
Mikrofone	42
Kopfhörer	44
Sonifikation	45
Sonderanwendungen	45
7. Beginn, erster Versuch der neuen Konzeption	47
Physikalisch-chemisch	47
Katalysator	47
Fotografische Entwicklerflüssigkeit	47
Kontrastmittel	47
Genetisch	48
Enzym	48
Hormon	48
Immunsystem	49

Reflex.....	49
Neuronal	49
Gewürze, den Geschmack betreffend.....	49
Parfum, den Geruch betreffend	49
Drogen, die Psychologie betreffend	50
Hypnose, die Bewusstheit betreffend.....	50
Gesellschaftlich, mehrere Teilnehmer betreffend.....	50
Lernen	50
Technisch.....	51
Vernetzt.....	52
Literatur	53