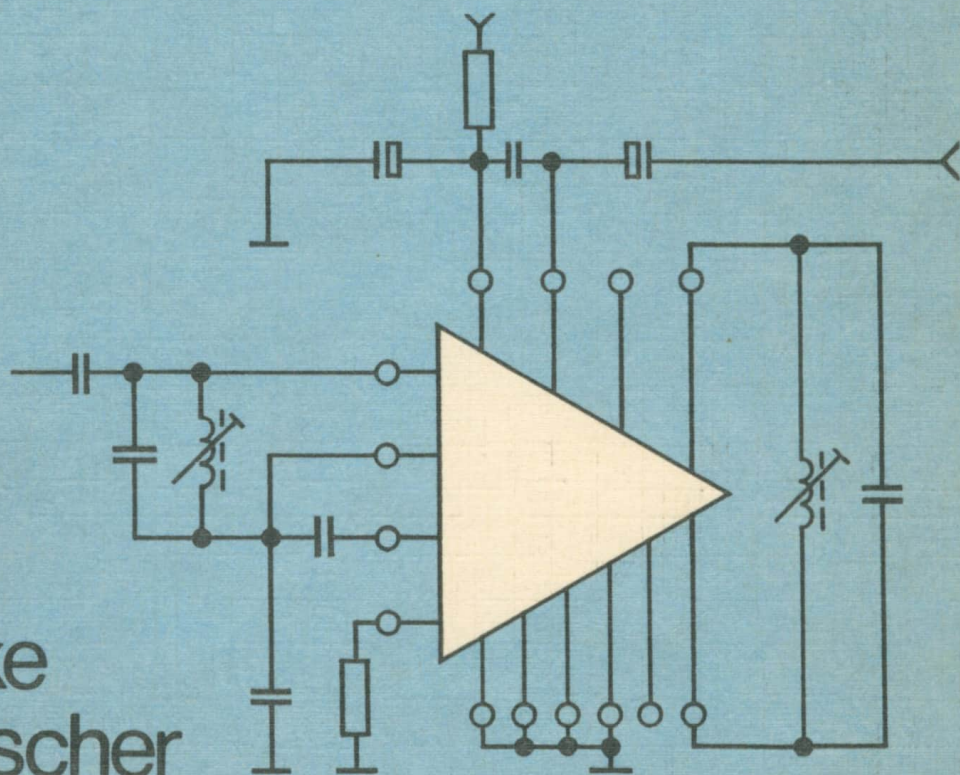


LEHRBUCH FÜR DIE BERUFSBILDUNG

Grund- schaltungen der Elektronik



Funke
Liebscher

LEHRBUCH FÜR DIE BERUFSBILDUNG

Grundsaltungen der Elektronik

Oberlehrer Dipl.-Gwl. Ing. Rainer Funke
Dipl.-Gwl. Dipl.-Ing. Siegfried Liebscher

13., bearbeitete Auflage



VEB VERLAG TECHNIK BERLIN

Als berufsbildende Literatur für die Ausbildung von Lehrlingen und Werk tätigen zum Facharbeiter
für verbindlich erklärt

1. 9. 1988

Ministerium für Elektrotechnik und Elektronik

Funke, Rainer:
Grundsaltungen der Elektronik / Rainer Funke :
Siegfried Liebscher. – 13., bearb. Aufl. – Berlin :
Verl. Technik, 1988 . – 192 S. : 367 Bilder,
14 Taf. – (Lehrbuch für die Berufsbildung)
ISBN 3-341-00564-1
NE: Liebscher, Siegfried

ISBN 3-341-00564-1

13., bearbeitete Auflage
© VEB Verlag Technik, Berlin, 1988
Lizenz 201 · 370/177/88
Printed in the German Democratic Republic
Gesamtherstellung: Nationales Druckhaus Berlin
Lektor: OL Dipl.-Gwl. Wolfgang Wosnizok
Einband: Kurt Beckert
LSV 3502 · VT 5/4255-13
Bestellnummer: 553 959 5
00850

Vorwort

Das vorliegende Lehrbuch hat die Aufgabe, ein erweiterungsfähiges Grundwissen so zu vermitteln, daß das innerhalb der Berufsausbildung in den jeweiligen Berufen erforderliche Spezialwissen, auf dem Grundwissen aufbauend, leicht erworben werden kann.

Die rasche Entwicklung der Elektronik machte eine abermalige Modernisierung des Inhalts für die 13. Auflage erforderlich. So wurden moderne Varianten der Stromversorgung eingearbeitet, die Ausführungen zum Operationsverstärker erweitert und dem heutigen Stand der Technik angepaßt, neue digitale Schaltkreisfamilien aufgenommen und der Umfang der digitalen Schaltungen vergrößert. Neu eingefügt wurde der Komplex AD/DA-Wandlung.

Durch die Darstellung des Stoffes in vielfältiger Form soll der Facharbeiter zu einer individuellen Weiterbildung für seine spätere berufliche Tätigkeit befähigt werden. So wurden grafische

und mathematische Lösungsverfahren zur Dimensionierung von Schaltungen an den jeweils überschaubaren Beispielen beschrieben sowie die Wirkungsweise von Grundsaltungen mathematisch, physikalisch oder grafisch erläutert. Entsprechend den unterschiedlichen Bildungsanforderungen in den einzelnen Elektronikberufen, sind neben einer ausführlicheren Darstellung von Schwerpunkten auch kürzer gefaßte Hinweise zu spezielleren Problemen enthalten. Dem Charakter und dem Umfang des Buches entsprechend, konnten keine dimensionierten Beispiele aufgenommen werden.

Wir wünschen allen Lesern, daß unser Buch die Aneignung eines ausbaufähigen Grundwissens erleichtern möge. Für Anregungen zur Verbesserung des Lehrbuchs sind wir stets dankbar. Unser Dank gilt allen, die am Gelingen der Auflage beteiligt waren.

Rainer Funke, Siegfried Liebscher

Inhaltsverzeichnis

Zur Schreibweise der verwendeten Größen	6	3.4.1. Allgemeines	53
1. Stromversorgungsschaltungen	7	3.4.2. Idealer Operationsverstärker	55
1.1. Allgemeines	7	3.4.3. Realer Operationsverstärker	55
1.2. Gleichrichterschaltungen	7	3.4.4. Grundsaltungen mit Operationsverstärkern	56
1.3. Schaltungen zur Glättung der pulsierenden Gleichspannung	10	3.4.5. Frequenzgang und seine Kompensation	59
1.3.1. Schaltung mit Ladekondensator	10	3.4.6. Nullpunktfehler und seine Kompensation	61
1.3.2. Siebschaltungen	11	3.4.7. Kenngrößen von Operationsverstärkern	62
1.4. Stabilisierungsschaltungen	15	3.4.8. Eigenschaften und Schaltungstechnik verschiedener Operationsverstärkertypen	64
1.4.1. Allgemeines	15	3.4.9. Operationsverstärker mit Stromdifferenzeingang (Norton-Verstärker)	70
1.4.2. Stabilisierungsschaltung mit Z-Diode	15	3.5. Vorverstärker	71
1.4.3. Elektronische Stabilisierungsschaltungen	17	3.5.1. Einstellung des Arbeitspunktes	71
1.5. Transverter	26	3.5.2. RC-gekoppelter Verstärker	73
1.5.1. Allgemeines	26	3.5.3. Bandfiltergekoppelter Verstärker	76
1.5.2. Sperrwandler	27	3.5.4. Direktgekoppelter Verstärker	77
1.5.3. Summierwandler	30	3.5.5. Optoelektronisch gekoppelter Verstärker	79
1.5.4. Stromflußwandler	30	3.6. Endverstärker	80
1.6. Schaltnetzteile	31	3.6.1. Allgemeines	80
1.7. Sonderprobleme bei Stromversorgungsschaltungen im Schaltbetrieb	32	3.6.2. Eintaktverstärker	81
1.8. Aufgaben	33	3.6.3. Gegentaktverstärker	82
2. Schaltungen zur Verminderung der Induktionsspannung bei plötzlichen Stromänderungen in Spulen	35	3.6.4. Endverstärker mit integrierten Schaltkreisen	89
2.1. Allgemeines	35	3.6.5. Endverstärker mit sehr hohem Wirkungsgrad (D-Verstärker)	93
2.2. Schaltungsvarianten	35	3.7. Einstellglieder im Niederfrequenzverstärker	95
2.3. Aufgaben	36	3.7.1. Verstärkungseinsteller	95
3. Verstärker mit elektronischen Bauelementen	37	3.7.2. Klangeinsteller	95
3.1. Verstärker als Vierpol	37	3.7.3. Elektronische Einstellglieder	96
3.2. Grundsaltungen der Verstärker	39	3.8. Aufgaben	97
3.2.1. Allgemeines	39	4. Sinus- und Kippgeneratoren	99
3.2.2. Grundsaltungen mit einem Transistor	39	4.1. Allgemeines	99
3.2.3. Grundsaltungen mit mehreren Transistoren	41	4.2. Sinusgeneratoren	100
3.3. Gegenkopplung	49	4.2.1. Meißner-Oszillatoren	100
3.3.1. Allgemeines	49	4.2.2. Dreipunktoszillatoren	101
3.3.2. Ersatzschaltungen des Verstärkerbauelements	49	4.2.3. Hochfrequenzoszillator in Basis-schaltung	102
3.3.3. Gegenkopplungsschaltungen	51	4.2.4. Quarzoszillatoren	102
3.4. Operationsverstärker	53	4.2.5. RC-Oszillatoren	103

4.2.6. Oszillatoren mit integrierten Schaltkreisen	105	5.5.3. Anschluß systemfremder Lasten	145
4.3. Kippgeneratoren	107	5.5.4. Dual- und Dezimalzähler	146
4.3.1. Allgemeines	107	5.5.5. Frequenzteiler	149
4.3.2. Rechteckspannungsgeneratoren	109	5.5.6. Schieberegister und Schiebering	150
4.3.3. Sägezahngeneratoren	111	5.5.7. BCD-zu-7-Segment-Dekoder	151
4.4. Aufgaben	115	5.5.8. Multiplexer	152
5. Elektronische Grundsaltungen der Digitaltechnik	116	5.6. Analog-Digital-Wandler und Digital-Analog-Wandler	154
5.1. Allgemeines	116	5.6.1. Allgemeines	154
5.2. Transistor als elektronischer Schalter	117	5.6.2. Verfahren zur Analog-Digital-Wandlung	156
5.3. Kombinatorische Schaltungen	118	5.6.3. Verfahren zur Digital-Analog-Wandlung	159
5.3.1. Allgemeines	118	5.7. Aufgaben	161
5.3.2. Grundgatter	119	6. Modulation und Demodulation	163
5.3.3. Interface-Schaltungen	131	6.1. Theoretische Grundlagen	163
5.4. Sequentielle Schaltungen	133	6.1.1. Allgemeines	163
5.4.1. Allgemeines	133	6.1.2. Darstellung der Modulationsarten	164
5.4.2. RS-Flip-Flop	133	6.2. Modulations- und Demodulationsverfahren und -schaltungen	173
5.4.3. JK-Flip-Flop	134	6.2.1. Amplitudenmodulation und -demodulation	173
5.4.4. D-Flip-Flop	136	6.2.2. Winkelmodulation und -demodulation	180
5.4.5. Signalgeneratoren	136	6.2.3. Pulsmodulation und -demodulation	187
5.4.6. Monostabile Elemente	139	6.3. Aufgaben	189
5.4.7. Schwellwertelemente	141	Sachwörterverzeichnis	191
5.5. Ausgewählte Teilschaltungen digitaler Geräte	143		
5.5.1. Signaleingabe über mechanische Kontakte	143		
5.5.2. Impulsverzögerung und Flankenversteilerung	144		

Zur Schreibweise der verwendeten Größen

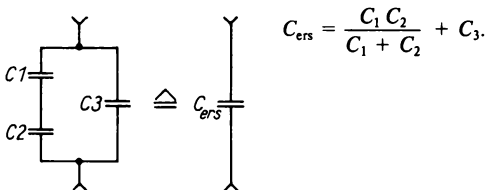
Bauelemente und Bauteile

Bauelemente und Bauteile werden mit standardisierten Kodebuchstaben bezeichnet. Gleichartige Bauelemente/Bauteile werden durch laufende Nummern hinter den Kodebuchstaben unterschieden (z. B. $C1$, $C2$). Erfolgt die Unterscheidung mit Buchstaben, die auf die Funktion in der Schaltung hinweisen, dann stehen diese als Index (z. B. C_L Ladekondensator, C_S Siebkondensator).

Gleichartige Schaltungsgrößen der Bauelemente (Kapazitäten, Widerstandswerte, Induktivitäten) werden durch Zahlen unterschieden, die als Index an den entsprechenden Formelzeichen stehen.

Beispiel

Die dargestellte Zusammenschaltung der Kondensatoren $C1$, $C2$ und $C3$ mit den Kapazitäten C_1 , C_2 und C_3 ergibt eine Ersatzkapazität C_{ers} von



Physikalische Größen

- Gleichartige physikalische Größen werden immer durch ein- oder mehrstellige Indizes unterschieden. Bei gerichteten Größen läßt sich damit auch eindeutig die Richtung angeben.

Beispiel

Zwischen den Punkten C und E liegt eine Spannung von 6 V.

C ist gegenüber E positiv.

Daraus resultiert, daß C im Index an erster Stelle, E an zweiter Stelle geschrieben wird:

$U_{CE} = 6 \text{ V}$ oder $U_{CE} = -6 \text{ V}$.

- Werden mehrere Spannungen auf ein gemeinsames, in Schaltungen durch ein Massezeichen gekennzeichnetes Potential 0 bezogen, dann wird die an zweiter Stelle im Index zu schreibende 0 weggelassen:

U_x anstatt U_{x0} . Der Punkt x ist positiv gegenüber Masse, z. B. $U_2 = 10 \text{ V}$.

- Eingangsgrößen erhalten bevorzugt den Index 1, Ausgangsgrößen den Index 2.
- Zeitunabhängige Ströme und Spannungen (Gleichströme und -spannungen) werden durch Großbuchstaben gekennzeichnet, z. B. I und U .
- Für zeitabhängige Ströme und Spannungen werden Kleinbuchstaben verwendet, z. B. i und u .
- Ist die Zeitabhängigkeit sinusförmig, steht als Index das Zeichen $\sim$, z. B. $I_{\sim}$ und $U_{\sim}$.
- Die arbeitspunktabhängigen Parameter der Bauelemente werden klein, alle anderen groß geschrieben.

Beispiel

r_1 Eingangswiderstand des Transistors
 R_{ein} Eingangswiderstand der Transistorschaltung.

- Die Kenngrößen rückgekoppelter Schaltungen werden, im Gegensatz zu nichtrückgekoppelten, mit einem hochgestellten Strich versehen.

Beispiel

Verstärkungsfaktor eines gegengekoppelten Verstärkers.

$$v' = \frac{U_2'}{U_1'}$$

Eine Ausnahme bilden die Kenngrößen bei Operationsverstärkerschaltungen. Da OV-Schaltungen stets gegengekoppelt sind, wird auf den hochgestellten Strich verzichtet.

- Leistungs-, Strom- und Spannungsverhältnisse werden häufig als logarithmische Verhältnisse in dB angegeben:

$$v = 10 \lg \frac{P_2}{P_1} \text{ dB} = 20 \lg \frac{I_2}{I_1} \text{ dB} = 20 \lg \frac{U_2}{U_1} \text{ dB}.$$

Beispiel

Das lineare Spannungsverhältnis $v = 0,5$ entspricht dem logarithmischen Spannungsverhältnis $v = -6 \text{ dB}$; das lineare Spannungsverhältnis $v = 2$ entspricht dem logarithmischen $v = 6 \text{ dB}$.

1.

Stromversorgungsschaltungen

1.1. Allgemeines

Fast alle elektronischen Geräte benötigen für ihre Funktion Gleichspannungen. Um diese Geräte an das Wechselstromnetz anschließen zu können, ist ein Stromversorgungsteil notwendig. In ihm wird der Wechselstrom gleichgerichtet und der pulsierende Gleichstrom geglättet. Stromversorgungsteile sind auch dann notwendig, wenn elektronische Geräte aus einer Gleichspannungsquelle (z. B. Autobatterie) gespeist werden sollen, deren Gleichspannung von der für das Gerät erforderlichen Gleichspannung abweicht. Bei höheren Anforderungen an die Konstanz der gewonnenen Gleichspannung wird diese stabilisiert. Versorgungsspannungsschwankungen oder Änderungen der Stromaufnahme des Gerätes beeinflussen dann nicht mehr die Höhe der bereitgestellten Gleichspannung.

1.2. Gleichrichterschaltungen

Bei einer Gleichrichtung wird eine Wechselspannung in eine Gleichspannung gewandelt. Dies kann mit steuerbarem Gleichrichter (Thyristor) oder nichtsteuerbarem Gleichrichter (Halbleitergleichrichter) erfolgen. In diesem Lehrbuch wird nur die zweite Art behandelt. Je nachdem, ob nur eine Halbperiode oder beide Halbperioden der Wechselspannung zur Gleichrichtung genutzt werden, spricht man von *Einweg-* oder *Zweiweggleichrichtung* (Bild 1.1). Bild 1.2 zeigt eine Einweggleichrichterschaltung (Reihenschaltung von Diode und Lastwiderstand). Über dem Widerstand liegt eine *pulsierende Gleichspannung* gemäß Bild 1.1 a. Um auch die zweite Halbperiode der Wechselspannung für die Gleichspannung auszunutzen, wendet man die Zweiweggleichrichterschaltung an. Man unterscheidet zwei Arten von Zweiweggleichrichterschaltungen, die *Gegentaktschaltung* und die *Brückenschaltung*.

Bild 1.3 zeigt die *Gegentaktschaltung* mit zwei Gleichrichterelementen und einem Transformator mit Mittelabgriff. Für beide Halbwellen der Wechselspannung sind die jeweils geöffneten Stromkreise dargestellt. Für jede Halbperiode der Wechselspannung ist jeweils eine Diode geöffnet. Die Ströme beider Halbwellen fließen gleichsinnig durch den Lastwiderstand. Die zweite Zweiweggleichrichterschaltung, die *Brückenschaltung* (Grätz-Schaltung), wird im Bild 1.4 gezeigt. Auch hier ist für beide Halb-

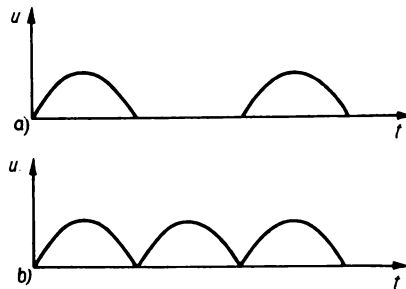


Bild 1.1. Spannungsverlauf

- a) nach Einweggleichrichtung;
- b) nach Zweiweggleichrichtung

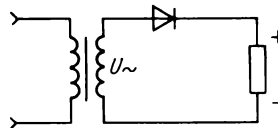


Bild 1.2. Einweggleichrichterschaltung

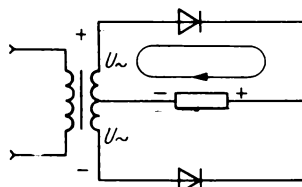


Bild 1.3. Zweiweggleichrichterschaltung (Gegentaktschaltung)

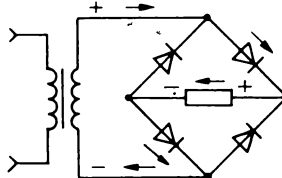


Bild 1.4. Zweiweggleichrichterschaltung (Brückenschaltung)

wellen der jeweilige Stromkreis gekennzeichnet. Es liegen immer zwei Gleichrichter in Reihe. Der eingangsseitig dargestellte Transformator ist bei einer Gegentaktschaltung immer erforderlich. Bei anderen Schaltungsvarianten ist er nur dann erforderlich, wenn eine galvanische Trennung vom Netz erreicht werden soll oder die Höhe der gewünschten Gleichspannung von der Höhe der speisenden Wechselspannung stark abweicht.

Sollte bei der Gegentaktschaltung der Lastkreis unterbrochen sein, so liegt an jedem Gleichrichterelement jeweils die Summe der Spitzenwerte der an den beiden Teilwicklungen liegenden Spannungen ($2 \hat{U}$) als Sperrspannung. Bei der Brückenschaltung tritt unter derselben Bedingung an jedem Gleichrichterelement der Spitzenwert der anliegenden Wechselspannung ($\hat{U}$) als Sperrspannung auf. Die Spannungsbelastung der Gleichrichter ist, bezogen auf gleiche Ausgangsspannung, bei der Gegentaktschaltung deshalb doppelt so hoch wie bei der Brückenschaltung, wodurch die Auswahl der Schaltung wesentlich bestimmt wird. Die Frequenz der pulsierenden Gleichspannung nach Zweweggleichrichtung wird doppelt so hoch wie die Frequenz der Wechselspannung.

Die Welligkeit der entstandenen Gleichspannung ist auch nach Zweweggleichrichtung noch sehr groß, weshalb ein erheblicher Aufwand an Siebmitteln erforderlich ist (s. Abschn. 1.3.). Werden große Gleichstromleistungen gefordert, so wendet man meist Drehstrom-Gleichrichterschaltungen an. Eine entsprechende Einwegschaltung wird im Bild 1.5 gezeigt. Der aus dem Drehstrom-Liniendiagramm abgeleitete Verlauf der gleichgerichteten Spannung ist im Bild 1.6 dargestellt. Eine Gleichspannung noch geringerer Welligkeit entsteht bei einer Drehstrom-Zweweggleichrichterschaltung. Auch hier sind sowohl eine Gegentaktschaltung (Bild 1.7) als auch eine Brückenschaltung (Bild 1.8) möglich. Das Drehstrom-Liniendiagramm wird im Bild 1.9 gezeigt. Es ist gekennzeichnet, welche der sechs Dioden der Brückenschaltung jeweils stromführend ist. Die Welligkeit der so entstandenen Gleichspannung ist wesentlich geringer als bei Einphasengleichrichterschaltungen.

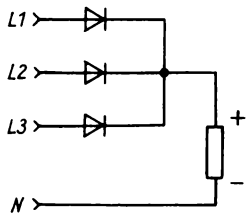


Bild 1.5. Drehstrom-Einweggleichrichterschaltung

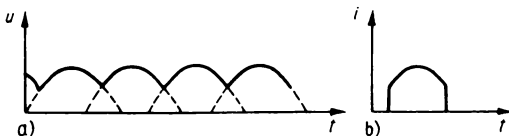


Bild 1.6. Liniendiagramm bei Drehstromgleichrichtung

a) Spannungsverlauf am Lastwiderstand; b) Stromverlauf durch eine Diode

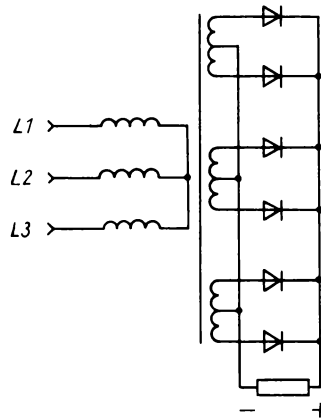


Bild 1.7. Drehstrom-Gegentaktgleichrichterschaltung

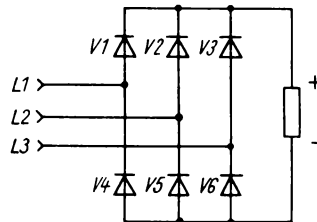


Bild 1.8. Drehstrom-Brückengleichrichterschaltung

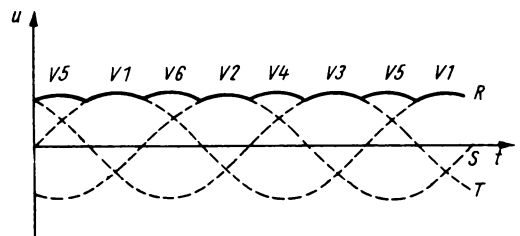


Bild 1.9. Liniendiagramm der Spannung nach Drehstrom-Zweweggleichrichtung

Mit der im Bild 1.10 dargestellten Spannungsverdopplerschaltung ist es möglich, bei Einphasengleichrichtung eine gegenüber der Wechselspannung höhere Gleichspannung zu erzielen. Während der positiven Halbwelle wird der Kondensator $C1$ über die Diode $V1$ etwa auf den Spitzenwert der Wechselspannung aufgeladen (schwarze Strompfeile), während der negativen Halbwelle entsprechend $C2$ über $V2$ (rote Strompfeile). Am Lastwiderstand liegt die Summe der Ladespannungen beider Kondensatoren. Diese müssen so bemessen sein, daß sie sich während einer Periode nur sehr wenig über den Lastwiderstand entladen. (Die Entladezeitkonstante muß wesentlich größer als die Periodendauer der Wechselspannung sein.)

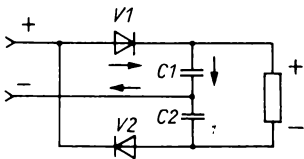


Bild 1.10. Spannungsverdopplerschaltung

Eine Spannungsverdopplung läßt sich auch erreichen, wenn man eine Schaltung nach Bild 1.11 aufbaut. Diese Verdopplerschaltung läßt sich durch mehrmaliges Hintereinanderschalten zu einer Vervielfacherschaltung erweitern (Bild 1.12). Liegt am oberen Eingang (Bild 1.11) die positive Halbwelle der Wechselspannung, so wird $C1$ über $V1$ in der eingetragenen

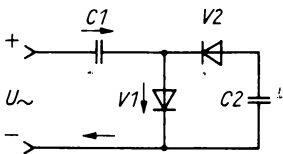


Bild 1.11. Spannungsverdopplerschaltung als Grundstufe der Spannungsvervielfacherschaltung

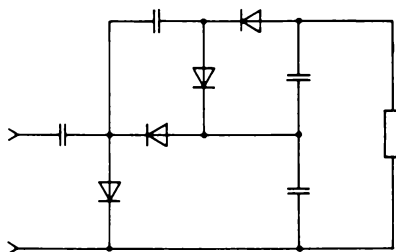


Bild 1.12. Spannungsvervielfacherschaltung

nen Polarität aufgeladen. $V2$ ist während dieser Zeit gesperrt. Bei der folgenden negativen Halbwelle am oberen Eingang wird $C2$ über $V2$ auf die doppelte Spitzenspannung der anliegenden Wechselspannung aufgeladen, da sich jetzt die Ladespannung von $C1$ zur angelegten Wechselspannung addiert.

In Geräten werden bevorzugt Siliziumgleichrichter eingesetzt. Ihr Vorteil ist ein hoher Wirkungsgrad und ein geringes Volumen. Beim praktischen Einsatz müssen einige Besonderheiten beachtet werden. Infolge ihrer geringen Masse sind sie bereits gegen kurzzeitige Überlastungen sehr empfindlich. Eine stoßweise Zufuhr der Wärmemenge Q muß nach

$$Q = m c \Delta t \quad (1.1)$$

eine höhere Temperaturzunahme bewirken als bei einem Selengleichrichter mit einer vergleichsweise viel höheren Masse. Der bei Überschreiten der Durchbruchspannung gegen 0 strebende Sperrwiderstand ist bei der Spannungsdimensionierung besonders zu beachten. Es treten im Wechselspannungsnetz Überspannungsspitzen bis zur vielfachen Höhe der Nennspannung mit einer Zeitdauer von einigen Millisekunden auf. Hervorgerufen werden sie durch Schaltvorgänge im Netz. Um Schäden an den Siliziumgleichrichtern zu vermeiden, muß man diese Überspannung wirkungsvoll begrenzen oder eine entsprechende spannungsmäßige Überdimensionierung vornehmen.

Die schädliche Wirkung des Trägerstau-effekts wird vermieden, wenn man parallel zu den Dioden die Reihenschaltung eines Kondensators mit einem Widerstand oder nur einen Kondensator schaltet (Bild 1.13). Gleichzeitig werden Überspannungsspitzen anderer Ursachen wirkungsvoll begrenzt.

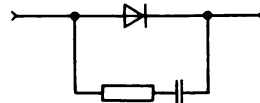


Bild 1.13. Schutzschaltung bei Siliziumdioden

Der Trägerstau-effekt entsteht beim Übergang von der Durchlaß- zur Sperrrichtung, da die Ladungsträger erst von der Sperrschicht abgezogen werden müssen. Hierzu ist eine bestimmte Zeit erforderlich, während der ein großer Sperrstrom fließen kann. Nach Ablauf dieser Zeit erfolgt der Übergang abrupt, was bei Schaltungen mit Induktivitäten zu Überspannungen führt. Ein weiches Umschalten gestatten Dioden mit soft recovery Charakteristik.

Zusammenfassung

Die Gleichrichtung von Wechselspannung ist zur Stromversorgung in elektronischen Geräten erforderlich. Die Auswahl der Gleichrichterschaltung richtet sich nach der zulässigen Welligkeit und dem ökonomisch vertretbaren Bauelementeaufwand. Bei der Einwegschaltung wird eine Halbwellen der Wechselspannung, bei der Zweiwegschaltung werden beide Halbwellen ausgenutzt. Eine gegenüber der Wechselspannung höhere Gleichspannung ermöglicht die Spannungsverdopplerschaltung.

Für hohe Gleichstromleistungen sind Drehstrom-Gleichrichterschaltungen zu bevorzugen. Bei Einsatz von Siliziumdioden sind deren besondere Eigenschaften (geringe Wärmeträgheit, hohe Überspannungsempfindlichkeit) besonders zu beachten.

1.3.

Schaltungen zur Glättung der pulsierenden Gleichspannung

1.3.1.

Schaltung mit Ladekondensator

Eine einfache Möglichkeit zur Glättung der pulsierenden Gleichspannung stellt der Einsatz eines Ladekondensators dar (Bild 1.14). Er wird während der ersten Halbwellen der pulsierenden Gleichspannung aufgeladen und gibt beim Sinken der Gleichspannung seine Ladung an den Lastwiderstand ab (Bild 1.15). Durch die Diode fließt nur dann ein Strom, wenn

$$u_D = U_{\sim} - u_C > 0 \text{ V} \quad (1.2)$$

ist. Die Zeitdauer des Stromflusses ist im Bild 1.15 durch den *doppelten Stromflußwinkel* gekennzeichnet. Der Stromflußwinkel Θ ist der halbe Winkel des Stromflusses, bezogen auf die ganze Periode 2π . Während der Sperrzeit entlädt sich der Kondensator über den Lastwiderstand R_L . Je größer die Entladezeitkonstante

$$\tau_{\text{ent}} = C R_L \quad (1.3)$$

ist, um so besser wird die Glättung der Gleichspannung sein. Eine höher werdende Kapazität des Kondensators verringert den Stromflußwinkel Θ , weshalb die Ladestromspitzen sehr hohe Werte annehmen. Sie sind dabei nur durch den Innenwiderstand der Gleichrichterschaltung begrenzt.

$$R_i = R_{Tr} + R_D + R_V \quad (1.4)$$

$R_{Tr} = R_2 + \ddot{u}^2 R_1$ (Transformatorwiderstand)

R_D Durchlaßwiderstand der Diode

R_V eventuell vorhandener zusätzlicher Vorwiderstand.

Der Gleichrichter wird im Einschaltmoment durch die Summe aus Verbraucher- und Ladestrom belastet. Unmittelbar nach dem Einschalten bedingt der noch ungeladene Kondensator einen sehr hohen Ladestromstoß. Dadurch wird die maximale Kapazität des Ladekondensators begrenzt.

Die auftretende Sperrspannung am Gleichrichter wird durch den Kondensator ebenfalls erhöht (Bild 1.15). Nimmt man dabei für den ungünstigsten Fall eine Aufladung des Kondensators auf den Spitzenwert der pulsierenden

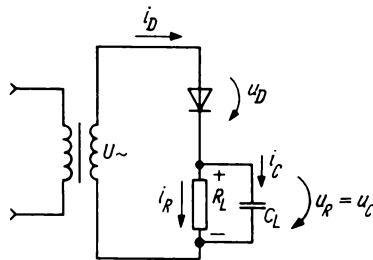


Bild 1.14. Gleichrichterschaltung mit Ladekondensator

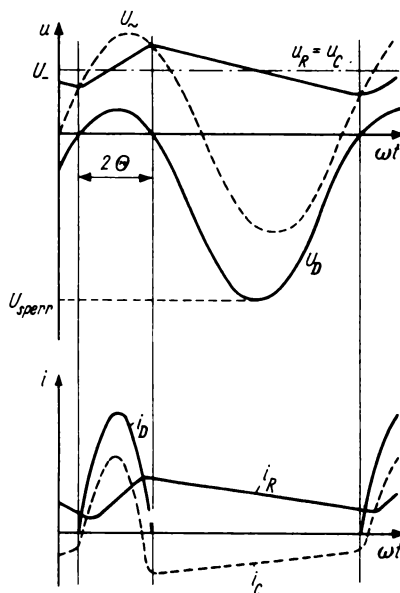


Bild 1.15. Darstellung der Strom- und Spannungsverläufe einer Gleichrichterschaltung mit Ladekondensator

Gleichspannung an, so tritt am Gleichrichter eine maximale Spannung in Sperrichtung von

$$u_{D \max} = 2 \sqrt{2} U_{\sim} \quad (1.5)$$

auf.

Eine genaue Berechnung der Gleichspannung am Ladekondensator und der verbleibenden Welligkeit ist schwierig. In der Praxis bedient man sich vielfach der in der Fachliteratur angegebenen Diagramme. Hier soll nur eine Näherung zur Berechnung des Wechselspannungsanteils der Gleichspannung am Ladekondensator angegeben werden. (Dieser Wechselspannungsanteil wird, wie auch in den nachfolgenden Betrachtungen, zur Vereinfachung als sinusförmig angesehen.)

$$U_{\sim} \approx \frac{U_-}{4\sqrt{2} p f R_L C_L} \quad (1.6)$$

U_- Gleichspannung am Ladekondensator
 p Anzahl der Gleichrichterwege
 f Frequenz der Wechselspannung
 R_L Lastwiderstand
 C_L Ladekondensator.

Beispiel

Eine Zweiweggleichrichterschaltung für 20 V wird mit einem Widerstand von 200 Ω belastet. Der Ladekondensator hat eine Kapazität von 500 μF . Wie groß ist die verbleibende Wechselspannung?

Gegeben: Gesucht:

$U_- = 20 \text{ V}$ $U_{\sim}$
 $p = 2$
 $f = 50 \text{ Hz}$
 $R_L = 200 \Omega$
 $C_L = 500 \mu\text{F}$

Lösung:

$$U_{\sim} \approx \frac{U_-}{4\sqrt{2} p f R_L C_L}$$

$$U_{\sim} \approx \frac{20 \text{ V}}{4\sqrt{2} \cdot 2 \cdot 50 \text{ s}^{-1} \cdot 2 \cdot 10^2 \Omega \cdot 5 \cdot 10^{-4} \text{ A} \cdot \text{s} \cdot \text{V}^{-1}}$$

$$U_{\sim} \approx 0,354 \text{ V}$$

Für die meisten praktischen Fälle ist diese verbleibende Wechselspannung zu hoch. Eine weitere Erniedrigung der Welligkeit erreicht man mit Siebschaltungen.

1.3.2.

Siebschaltungen

Siebschaltungen werden vorrangig als LC-Siebglied oder als RC-Siebglied ausgeführt

(Bild 1.16). Die Anordnung der Siebschaltung innerhalb einer Gleichrichterschaltung wird häufig als Siebkette bezeichnet (Bild 1.17). Bei der Wirkungsweise betrachtet man das Wechsel- und Gleichstromverhalten getrennt. In beiden Fällen handelt es sich um einen Spannungsteiler – für Wechselspannung aus der Induktivität L (bzw. Siebwiderstand R_S) und dem Siebkondensator C_S , für Gleichspannung aus dem Wicklungswiderstand der Induktivität (bzw. Siebwiderstand R_S) und dem Lastwiderstand R_L gebildet (Bild 1.18).

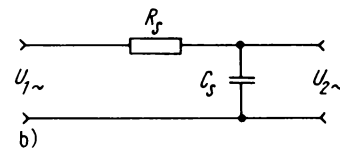
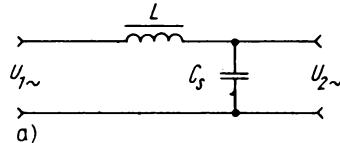


Bild 1.16. Siebglieder

a) LC-Siebglied; b) RC-Siebglied

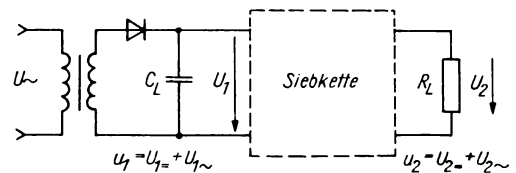


Bild 1.17. Anordnung der Siebkette innerhalb einer Gleichrichterschaltung

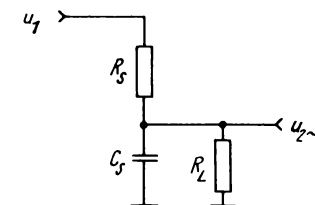
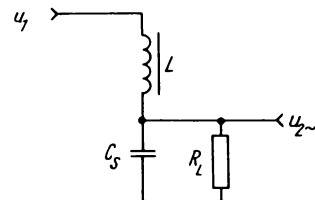


Bild 1.18. Darstellung der Siebglieder als Wechselspannungsteiler

Das Wechselspannungsverhalten wird durch den Siebfaktor S gekennzeichnet:

$$S = \frac{U_{1\sim}}{U_{2\sim}} \quad (1.7)$$

Berechnung der Siebfaktoren

LC-Siebglied

Der oben beschriebene Wechselspannteiler darf nur dann zur Berechnung herangezogen werden, wenn der kapazitive Widerstand des Siebkondensators sehr viel kleiner als der Lastwiderstand ist und der Wicklungswiderstand der Induktivität gegenüber deren induktivem Widerstand vernachlässigbar ist.

Aus der Spannungsteilerregel folgt unter diesen Bedingungen:

$$S = \left| \frac{U_{1\sim}}{U_{2\sim}} \right| \approx \left| \frac{\sqrt{\left(\omega L - \frac{1}{\omega C_s}\right)^2}}{\frac{1}{\omega C_s}} \right|$$

$$S \approx \left| \frac{\frac{\omega L \omega C_s - 1}{\omega C_s}}{\frac{1}{\omega C_s}} \right|$$

$$S \approx |\omega^2 LC_s - 1|.$$

Für $\omega^2 LC_s \gg 1$ kann geschrieben werden

$$S \approx \omega^2 LC_s \quad (1.8)$$

Zur einfachen Bestimmung des Siebfaktors kann man ein Nomogramm nach Bild 1.19 be-

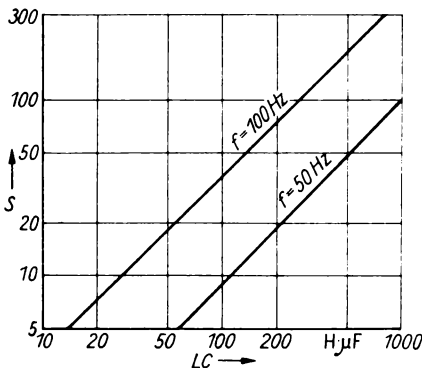


Bild 1.19. Nomogramm zur Bestimmung des Siebfaktors eines LC-Siebglieds

nutzen. Auf ihm ist der Siebfaktor in Abhängigkeit von dem Produkt aus L in H und C in μF dargestellt. Parameter dieser Darstellung ist die Frequenz der der Gleichspannung überlagerten Wechselspannung.

Beispiel

Wie groß ist die Restwechselspannung am Ausgang eines LC-Siebglieds mit $L = 10$ H und $C = 50$ μF hinter einer Einweggleichrichterschaltung, wenn die Wechselspannung am Ladekondensator 5 V beträgt?

Gegeben:

$$\begin{aligned} U_{1\sim} &= 5 \text{ V} \\ L &= 10 \text{ H} \\ C &= 50 \mu F \\ p &= 1 \end{aligned}$$

Gesucht:

$$U_{2\sim}$$

Lösung:

$$U_{2\sim} = \frac{U_{1\sim}}{S}, \quad S \approx \omega^2 LC$$

$$U_{2\sim} \approx \frac{U_{1\sim}}{\omega^2 LC}$$

$$U_{2\sim} \approx \frac{5 \text{ V}}{3,14^2 \cdot 10^4 \text{ s}^{-1} \cdot 10 \text{ V} \cdot \text{s} \cdot \text{A}^{-1} \cdot 50 \cdot 10^{-6} \text{ A} \cdot \text{s} \cdot \text{V}^{-1}}$$

$$U_{2\sim} \approx 102 \text{ mV}$$

Man erhält das gleiche Ergebnis, wenn man das Produkt aus L in H und C in μF bildet (hier 500), den Siebfaktor aus dem Nomogramm nimmt (hier 50) und dann $U_{2\sim}$ bestimmt:

$$U_{2\sim} \approx \frac{5 \text{ V}}{50} \quad U_{2\sim} \approx 100 \text{ mV}$$

RC-Siebglied

Die näherungsweise Berechnung erfolgt auch hier unter der Voraussetzung

$$X_{Cs} \ll R_L$$

$$S = \frac{U_{1\sim}}{U_{2\sim}} \approx \frac{\sqrt{R_s^2 + \left(\frac{1}{\omega C_s}\right)^2}}{\frac{1}{\omega C_s}},$$

$$S \approx \frac{\sqrt{R_s^2 \omega^2 C_s^2 + 1}}{\frac{1}{\omega C_s}}$$

$$S \approx \sqrt{R_s^2 \omega^2 C_s^2 + 1}$$

Für $R_s^2 \omega^2 C_s^2 \gg 1$ kann vereinfacht geschrieben werden:

$$S \approx \omega R_s C_s \quad (1.9)$$

Ein Nomogramm zur Bestimmung des Siebfaktors eines RC-Siebglieds ist im Bild 1.20 dargestellt.

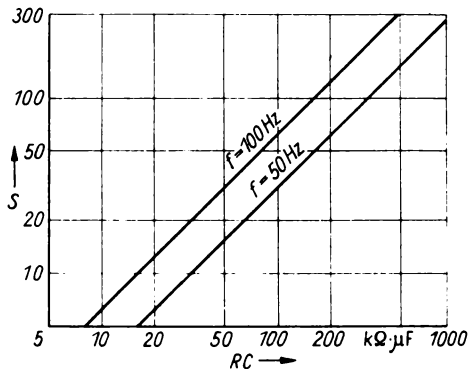


Bild 1.20. Nomogramm zur Bestimmung des Siebfaktors eines RC-Siebglieds

Beispiel

Wie groß ist der erforderliche Siebwiderstand eines RC-Siebglieds hinter einer Zweiweggleichrichterschaltung mit einem Siebkondensator von $50 \mu\text{F}$ bei einem geforderten Siebfaktor von 50?

Gegeben:

$C_s = 50 \mu\text{F}$
 $p = 2$
 $S = 50$

Gesucht:

R_s

Lösung:

Aus dem Nomogramm entnimmt man für $S = 50$ und $f = 100 \text{ Hz}$ das Produkt $R_{k\Omega} C_{\mu\text{F}}$ zu 80

$$R_{k\Omega} \approx \frac{80}{50}$$

$$R \approx 1,6 \text{ k}\Omega$$

An diesem Lösungsbeispiel erkennen wir, daß verhältnismäßig große Widerstandswerte erforderlich sind. Die am Siebwiderstand abfallende Gleichspannung und die damit verbundene große Verlustleistung ermöglichen den wirtschaftlichen Einsatz eines RC-Siebglieds nur bei kleinen Lastströmen (Vorverstärker, Mikrofonverstärker).

Beispiel

Dem RC-Siebglied des vorigen Lösungsbeispiels soll ein Gleichstrom von 3 mA entnommen werden. Wie groß sind der Spannungsabfall und die Verlustleistung am Siebwiderstand?

Gegeben:

$I_- = 3 \text{ mA}$
 $R = 1,6 \text{ k}\Omega$

Gesucht:

U_-
 P_v

Lösung:

$$U = I R$$

$$U = 3 \text{ mA} \cdot 1,6 \text{ k}\Omega$$

$$U = 4,8 \text{ V}$$

$$P_v = I^2 R$$

$$P_v = 9 \cdot 10^{-6} \text{ A}^2 \cdot 1,6 \cdot 10^3 \Omega$$

$$P_v = 14,4 \text{ mW}$$

Mehrgliedrige Siebkette

Erreicht man mit einem Ladekondensator und einem Siebglied keine ausreichend niedrige Welligkeit, so schaltet man mehrere Siebglieder zu einer mehrgliedrigen Siebkette hintereinander (Bild 1.21).

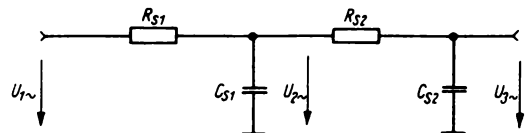


Bild 1.21. Mehrgliedrige RC-Siebkette

$$S = \frac{U_{1\sim}}{U_{3\sim}}$$

$$U_{3\sim} = \frac{U_{2\sim}}{S_2}$$

$$S = \frac{U_{1\sim} S_2}{U_{2\sim}}$$

$$\frac{U_{1\sim}}{U_{2\sim}} = S_1$$

$$S = S_1 S_2$$

$$(1.10)$$

Der Gesamtsiebfaktor ergibt sich als Produkt aus den Einzelsiebfaktoren, wenn die Voraussetzung $R_{S2} \gg X_{CS1}$ erfüllt ist. Mehrgliedrige RC-Siebketten wendet man gern zur Einsparung der teuren und voluminösen Siebdrossel in Rundfunk- und Fernsehgeräten an. Nach dem ersten Siebglied (mit niedrig bemessenem Siebwiderstand) entnimmt man den Strom für die Endstufe (höchster Laststrom, geringste Anforderung an die Welligkeit); am Ende des zweiten oder dritten Siebglieds mit hochohmigem Siebwiderstand erhält man eine Gleichspannung äußerst niedriger Welligkeit bei kleinstem Laststrom für brummempfindliche Vorverstärkerstufen.

Spezielle Siebschaltungen

Wie bereits erläutert, wird bei Einsatz eines Ladekondensators der Gleichrichter nur während einer kleinen Stromflußzeit vom Strom durchfließen. Während dieser kurzen Zeit ist der den Gleichrichter durchfließende Strom viel höher als der eigentliche Laststrom (s. Bild 1.15). Ein weiterer Nachteil des Ladekondensators ist die Erhöhung der Spannungsbelastung des Gleich-

richters. Die genannten Nachteile treten nur vermindert ein, wenn man den Einsatz eines Ladekondensators umgeht und dafür die Siebinduktivität mit nutzt (Bild 1.22). Da hier die Induktivität als Energiespeicher benutzt wird, sind ausreichend große Induktivitäten erforderlich. Die auftretende größere Brummspannung muß durch eine entsprechende Dimensionierung der Siebkette ausgeglichen werden.

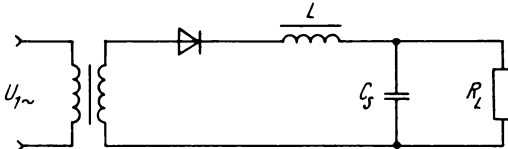


Bild 1.22. Siebschaltung ohne Ladekondensator

Teilweise werden auch Transistoren zur Siebung eingesetzt (Bild 1.23). Ökonomisch vertretbar wird diese Schaltung, wenn billige Leistungstransistoren zur Verfügung stehen. An sie werden keinerlei Anforderungen bezüglich ihrer Grenzfrequenz gestellt. Um die Wirkungsweise dieser Schaltung zu erläutern, wurde sie in eine Darstellung gemäß Bild 1.24 umgezeichnet. Man erkennt, daß über R_V der erforderliche Basisstrom zugeführt und in Verbindung mit R_2 und dem Kondensator das Basispotential weitgehend konstant gehalten wird. R_2 kann auch durch eine Z-Diode ersetzt werden. Der Lastwiderstand liegt in der Emittenerleitung; er verur-

sacht hier eine sehr starke Gleichstromgegenkopplung des Transistors. Liegt beispielsweise die positive Halbwelle der Brummspannung am Eingang, so tritt eine geringfügige Stromerhöhung durch den Lastwiderstand auf. Als Folge steigt das Emittenerpotential, und die wirksame Basis-Emittener-Spannung sinkt. Der Durchlaßwiderstand des Transistors und damit der Spannungsabfall an ihm werden erhöht. Der Siebfaktor dieser Schaltung hängt von dem Siebfaktor des RC -Gliedes (R_V, C) und den Transistordaten ab. Bei hoher Stromverstärkung des Transi-

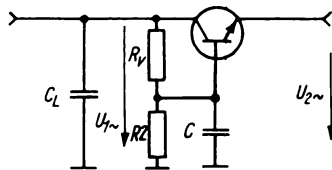


Bild 1.23. Einsatz eines Transistors zu Siebung

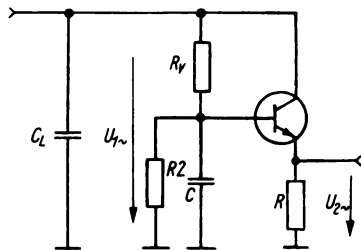


Bild 1.24. Zur Wirkungsweise der Siebung mit Transistor

Tafel 1.1. Siebschaltungen

Bezeichnung	Schaltung	Berechnung des Siebfaktors	Vor- und Nachteile
LC-Siebglied		$S \approx \omega^2 L C$	hoher Siebfaktor, Siebdrossel ist teuer und voluminös, verursacht magnetische Störfelder
RC-Siebglied		$S \approx \omega R C$	guter Siebfaktor bei geringem Aufwand, hoher Gleichspannungsabfall und hohe Verlustleistung am Siebwiderstand
Transistor-siebglied			sehr hoher Siebfaktor, Leistungstransistor erforderlich

stors kann der Siebfaktor des RC -Gliedes sehr groß gewählt werden (großes R), wodurch eine hohe Glättung des Emitterstroms erreicht wird. In Tafel 1.1 sind die Siebschaltungen übersichtlich gegenübergestellt.

Zusammenfassung

Die erforderliche Glättung der gleichgerichteten Spannung erfolgt meist mit Hilfe eines Ladekondensators und einem nachgeschalteten Siebglied. Die höchstmögliche Kapazität des Ladekondensators wird durch die auftretenden Ladestromspitzen begrenzt. Mit LC -Siebgliedern erreicht man hohe Siebfaktoren bei niedrigem Spannungs- und Leistungsverlust. Beim Verzicht auf die teure und große Siebdrossel entsteht beim RC -Siebglied der Nachteil hoher Spannungs- und Leistungsverluste. RC -Siebglieder sind deshalb nur bei niedrigen Lastströmen für hohe Siebfaktoren wirtschaftlich ausführbar. Häufig wendet man die Hintereinanderschaltung mehrerer RC -Siebglieder zu einer Siebkette an. Die Wirkung des Siebkondensators kann man durch Einsatz von Transistoren wesentlich erhöhen.

1.4. Stabilisierungsschaltungen

1.4.1. Allgemeines

Für viele elektronische Geräte wird eine konstante Betriebsspannung gefordert. Betriebsspannungsschwankungen treten als Folge von Netzspannungsschwankungen, Änderungen der Umgebungstemperatur oder Belastungsänderungen auf.

Eine Konstanzhaltung der Spannung ist mit steuerbaren Bauelementen oder mit Bauelementen, deren Strom-Spannungs-Kennlinie

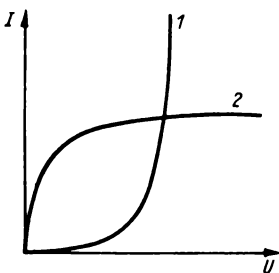


Bild 1.25. Kennlinien nichtlinearer Bauelemente

nichtlinear ist, möglich. Bauelemente mit geringem Wechselstromwiderstand (Bild 1.25, Kennlinie 1) eignen sich besonders zur Spannungsstabilisierung, solche mit hohem Wechselstromwiderstand (Bild 1.25, Kennlinie 2) zur Stromstabilisierung.

Der Grad der Stabilisierung wird durch den Stabilisierungsfaktor S_i beschrieben:

$$S_i = \frac{\frac{\Delta U_1}{U_1}}{\frac{\Delta U_2}{U_2}} \quad (1.11)$$

1.4.2. Stabilisierungsschaltung mit Z-Diode

Der nahezu senkrecht verlaufende Teil der Kennlinie einer Z-Diode wird zur Stabilisierung genutzt. Die stabilisierende Wirkung dieser Schaltung (Bild 1.26) wird in einer grafischen Darstellung des Stromkreises besonders anschaulich. Zunächst soll die Konstruktion für den lastseitigen Leerlauf erfolgen ($R_L \gg R_V$). Als passiven Zweipol betrachtet man das Bauelement mit der nichtlinearen U - I -Kennlinie, hier die Z-Diode. Der Widerstand R_V wird zusammen mit der Eingangsspannung U_1 als aktiver Zweipol aufgefaßt. Die entstehende Kennlinie dieses aktiven Zweipols wird häufig als Widerstandsgerade bezeichnet. Für zwei möglichen Eingangsspannungen ist die Konstruktion im Bild 1.27 dargestellt. Die Schnittpunkte der Kennlinien des aktiven Zweipols und der Z-Dioden-Kennlinie ergeben die beiden möglichen Spannungen an der Z-Diode und damit die mögliche Ausgangsspannungsschwankung bei

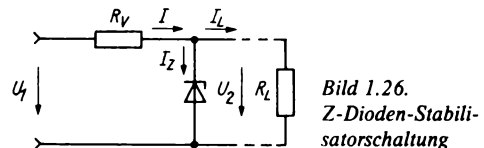


Bild 1.26. Z-Dioden-Stabilisatorschaltung

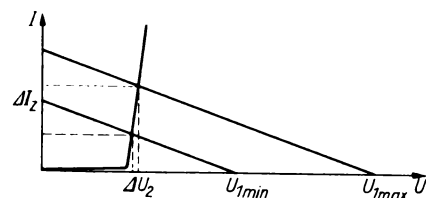


Bild 1.27. Darstellung der stabilisierenden Wirkung an der Kennlinie der Z-Diode

vorgegebener Eingangsspannungsschwankung. Gleichzeitig ist zu erkennen, daß die Änderung der Eingangsspannung zu einer großen Stromänderung ΔI_Z geführt hat. Die stabilisierende Wirkung beruht folglich auf der durch ΔI_Z hervorgerufenen Änderung des Spannungsabfalls am Vorwiderstand. Dieser muß dabei so bemessen sein, daß der Schnittpunkt seiner Geraden mit der Z-Dioden-Kennlinie stets im Arbeitsbereich derselben liegt.

Im folgenden wird zunächst ein *grafisches Lösungsverfahren* zur Bestimmung des erforderlichen Wertes des Vorwiderstands und der maximalen Spannungsänderung besprochen. Eine entsprechende Verfahrensweise ist auch bei vielen anderen Aufgaben der Elektronik anwendbar (z. B. Bestimmung der Widerstände zur Arbeitspunkteinstellung am Transistor). Im Normalfall ist die Z-Diode durch einen Lastwiderstand belastet. Es muß deshalb die resultierende Kennlinie dieser Parallelschaltung aus Z-Diode und Lastwiderstand konstruiert werden. Da sich bei einer Parallelschaltung für eine vorgegebene Spannung die Teilströme addieren, findet man zu jedem Spannungswert im Kennlinienfeld den zugehörigen Stromwert durch Addition der Stromwerte der Z-Dioden-Kennlinie und der Widerstandskennlinie des Lastwiderstands (Bild 1.28).

Man verfährt bei der Konstruktion so, daß man zunächst die Kennlinie der Z-Diode darstellt und dann die U - I -Kennlinie des Lastwiderstands einträgt. Für einige Spannungswerte bestimmt man grafisch den Strom durch den Lastwiderstand und die Z-Diode, wobei die Addition der Ströme bei der jeweiligen Spannung die Punkte für die resultierende Kennlinie der Parallelschaltung ergeben (im Bild 1.28 rot eingetragen). Die Widerstandsgerade des Vorwiderstands legt man so, daß bei der kleinsten Eingangsspannung die Widerstandsgerade um den minimalen Z-Strom oberhalb des Knickes der

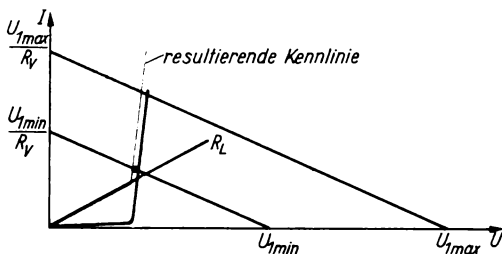


Bild 1.28. Grafische Bestimmung der erforderlichen Größe des Vorwiderstandes

resultierenden Kennlinie verläuft und nach Parallelverschiebung bei der größten Eingangsspannung das durch den maximal zulässigen Z-Dioden-Strom bestimmte obere Ende der resultierenden Kennlinie nicht überschritten wird. Soll die Schaltung auch bei Lastausfall sicher arbeiten, dürfte die Widerstandsgerade das obere Ende der Z-Dioden-Kennlinie nicht überschreiten. Wenn beide Bedingungen eingehalten sind, ist die Z-Diode für den vorgegebenen Fall einsetzbar. Den Wert des Vorwiderstands ermittelt man entweder aus der Steigung der Widerstandsgeraden oder aus ihrem Schnittpunkt mit der Stromachse.

Die Ausgangsspannungsschwankung hängt stark vom Verhältnis der Eingangsspannung zur Z-Dioden-Spannung ab. Man wählt für die Eingangsspannung etwa den 2- bis 4fachen Wert der zu stabilisierenden Spannung.

Beispiel

Es ist eine Spannung von 6,2 V bei einem Laststrom von 10 mA zu stabilisieren. Die Betriebsspannung betrage 20 V mit der relativen Schwankung $\Delta U/U = \pm 20\%$. Die Umgebungstemperatur der Z-Diode sei $t_a \leq 50^\circ\text{C}$. Es sind eine geeignete Z-Diode und der erforderliche Vorwiderstand zu bestimmen. Die Z-Diode soll auch bei Lastausfall nicht überlastet sein. Welcher Stabilisierungsfaktor wird erreicht?

Gegeben:

$$U_2 = 6,2 \text{ V}$$

$$I_L = 10 \text{ mA}$$

$$U_1 = 20 \text{ V} \pm 20\%$$

Gesucht:

geeignete Z-Diode

geeigneter Vorwiderstand

Ausgangsspannungsschwankung

Stabilisierungsfaktor

Graphische Lösung:

Aus dem Katalog entnimmt man als geeignete Z-Diode SZX 19/6,2 mit den für die Lösung benötigten Grenz- und Kennwerten: $r_z \leq 35 \Omega$; $P_{\text{tot}} = 500 \text{ mW}$ bei $t_a = 25^\circ\text{C}$; $R_{\text{th}} = 0,3 \text{ K/mW}$; $t_{\text{max}} = 175^\circ\text{C}$. Unter Vernachlässigung der Auslieferungstoleranz ist aus den Katalogangaben die Z-Dioden-Kennlinie darstellbar. Ihre Steigung wird durch r_z bestimmt. Der höchstzulässige Strom ergibt sich aus der zulässigen Verlustleistung unter den vorgegebenen Randbedingungen.

$$P_{\text{tot}} = \frac{t_j - t_a}{R_{\text{th}}} = \frac{(175 - 50) \text{ K} \cdot \text{mW}}{0,3 \text{ K}} = 416 \text{ mW}$$

$$I_{z\text{max}} = \frac{P_{\text{tot}}}{U_z} = \frac{416 \text{ mW}}{6,2 \text{ V}} = 67,2 \text{ mA}$$

Die resultierende Kennlinie ist im Bild 1.29 rot dargestellt. Zunächst werden die Grenzen für R_v bestimmt. Bei $R_{v\text{min}}$ darf im vorliegenden Fall bei der höchsten Betriebsspannung und Wegfall der Last der höchstzulässige Z-Diodenstrom nicht überschritten werden (obere gestrichelte Linie) und bei $R_{v\text{max}}$ sollte bei der

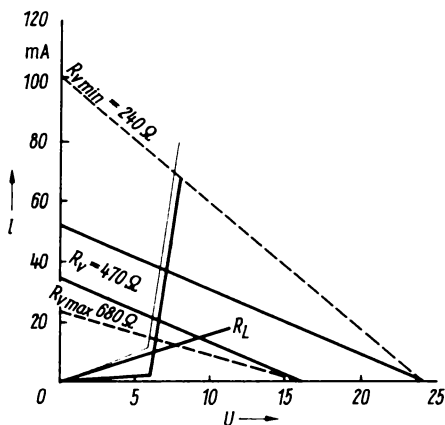


Bild 1.29. Konstruktive Ermittlung des Vorwiderstandes

kleinsten Betriebsspannung ein Mindeststrom durch die Z-Diode fließen (untere gestrichelte Linie). R_v muß somit in den Grenzen zwischen $240\ \Omega$ und $680\ \Omega$ liegen. Soll aus ökonomischen Gründen ein möglichst hoch tolerierter Widerstand (Reihe E6) verwendet werden, wäre ein Widerstand $470\ \Omega \pm 20\%$ einsetzbar (Toleranzbereich: $376\ \Omega \dots 564\ \Omega$). Für seinen Nennwert ist die Ausgangsspannungsschwankung graphisch ermittelt ($0,6\text{ V}$).

Die am Vorwiderstand auftretende Verlustleistung erreicht bei maximaler Eingangsspannung sowie kleinstmöglichem Wert von R_v ihren Höchstwert.

$$P_v = \frac{(U_{1\max} - U_z)^2}{R_{v\min}} \approx \frac{(24\text{ V} - 6,2\text{ V})^2}{376\ \Omega} = 0,84\text{ W} \quad (1.12)$$

Damit ist ein Widerstand einzusetzen, der auch bei 50°C mit mindestens $0,84\text{ W}$ belastbar ist (z. B. Widerstand der Kenngröße 25,518).

Der Stabilisierungsfaktor ergibt sich zu

$$S_i = \frac{\Delta U_1 / U_1}{\Delta U_2 / U_2} = \frac{8\text{ V} / 20\text{ V}}{0,6\text{ V} / 6,2\text{ V}} = 4,1$$

Rechnerische Lösung:

Für den Vorwiderstand betrachtet man dazu die bereits genannten Grenzbedingungen. [Zur Vereinfachung wird bei allen Berechnungen von einer konstanten Z-Dioden-Spannung ausgegangen (hier $6,2\text{ V}$). Die Ergebnisse sind entsprechend zu runden.] Für den größten zulässigen Vorwiderstand folgt

$$R_{v\max} = \frac{U_{1\min} - U_z}{I_L + I_{Z\min}} \approx \frac{16\text{ V} - 6,2\text{ V}}{10\text{ mA} + 5\text{ mA}} = 653\ \Omega. \quad (1.13)$$

Der kleinstzulässige Widerstand ergibt sich zu

$$R_{v\min} = \frac{U_{1\max} - U_z}{I_{Z\max}} \approx \frac{24\text{ V} - 6,2\text{ V}}{67,2\text{ mA}} = 264\ \Omega. \quad (1.14)$$

Mit diesen beiden Grenzwerten kann wiederum ein geeigneter Widerstand ausgewählt werden.

Die maximale Ausgangsspannungsschwankung bei angeschlossener Last ergibt sich zu

$$\Delta U_2 \approx \Delta I_z r_z \approx \left(\frac{U_{1\max} - U_z}{R_v} - \frac{U_{1\min} - U_z}{R_v} \right) r_z \quad (1.15)$$

$$\Delta U_2 \approx \frac{8\text{ V}}{470} \cdot 35\ \Omega \approx 0,59\text{ V}$$

Stabilisierungsfaktor: $S_i = 4,2$

Neben Z-Dioden werden zur Spannungsstabilisierung auch in Durchlaßrichtung betriebene Dioden verwendet, da diese nach Überschreiten der Schleusenspannung einen geringen differentiellen Widerstand haben. Mit einer Siliziumdiode läßt sich eine stabilisierte Spannung von $0,7\text{ V}$ erreichen, bei Reihenschaltung mehrerer Dioden ganzzahlige Vielfache von $0,7\text{ V}$.

1.4.3.

Elektronische Stabilisierungsschaltungen

Für die Betriebsspannungsversorgung vieler Geräte reicht die mit Z-Dioden erreichbare Stabilisierung der Gleichspannung nicht aus. Eine wesentlich bessere Konstanz der Ausgangsspannung (Erniedrigung des Innenwiderstands der Ersatzspannungsquelle) ist mit elektronischen Stabilisierungsschaltungen zu erreichen. Dabei können sowohl *stetige Regler* als auch *Schaltregler* eingesetzt werden.

Stetige Regler

Es sind zwei Schaltungsvarianten möglich:

- Der Stelltransistor wird als steuerbarer Vorwiderstand in Reihe mit dem Lastwiderstand (Reihenregelung) geschaltet (Bild 1.30).

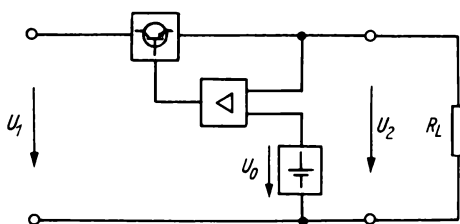


Bild 1.30. Prinzip der Reihenregelung

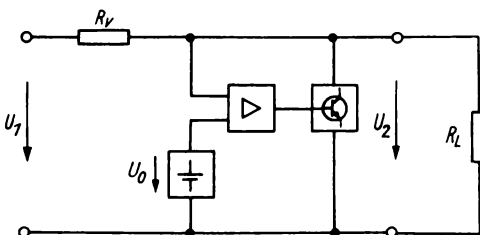


Bild 1.31. Prinzip der Parallelregelung

- Ein fester Vorwiderstand wird in Reihe zu dem parallel zum Lastwiderstand liegenden Stelltransistor geschaltet (Parallelregelung), der hier als steuerbarer Belastungswiderstand wirkt (Bild 1.31).

In beiden Schaltungen wird die Ausgangsspannung mit einer Vergleichsspannung U_0 (Referenzspannung) verglichen und die Abweichung vom Sollwert dem Verstärker zugeführt, der den Stelltransistor steuert.

Die Aufgabe des Stelltransistors der Reihenregelung ist aus seinem Kennlinienfeld (Bild 1.32) erkennbar. Es zeigt die Verhältnisse bei *konstantem Laststrom*, wobei die Widerstandsgerade des Lastwiderstands sowohl für die Nenn-eingangsspannung als auch für die minimale und die maximale Eingangsspannung dargestellt ist. Bei konstantem Laststrom wird der jeweilig erforderliche Arbeitspunkt durch die Schnittpunkte einer horizontalen Linie für den Nennlaststrom mit den einzelnen Widerstandsgeraden fixiert. Am Lastwiderstand liegt dann jeweils die gleiche Spannung an. Bei *konstanter Eingangsspannung* und veränderlichem Laststrom ergeben sich die Verhältnisse nach Bild 1.33. Soll die Ausgangsspannung trotz schwankender Last (gekennzeichnet durch drei Widerstandsgeraden mit unterschiedlicher Steigung) konstant bleiben, so muß im Abstand der Ausgangsspannung von der Betriebsspannung eine Vertikale eingetragen werden, deren Schnittpunkte mit den Widerstandsgeraden auch hier die erforderlichen Arbeitspunkte fixieren. Die selbsttätige Einstellung der geforderten Arbeitspunkte muß über den Regelverstärker am Stelltransistor erfolgen. Bild 1.34 zeigt eine komplette Schaltung. Die Vergleichsspannung wird durch eine Z-Diode erzeugt und bildet das Emittential des Verstärkertransistors. Die Ausgangsspannung wird über einen Spannungsteiler der Basis des Verstärkertransistors zugeführt. Dessen wirksame Basis-Emitter-Spannung ist folglich die Differenz zwischen Vergleichs- und Ausgangsspannung. Der Arbeitswiderstand des Verstärkertransistors bildet gleichzeitig den Basisvorwiderstand des Stelltransistors. Je nach Abweichung der Ausgangsspannung vom Sollwert wird sich ein bestimmter Spannungsabfall am Widerstand R_V einstellen, wodurch der Basisstrom des Stelltransistors gesteuert wird.

Wenn die Ausgangsspannung durch äußere Einflüsse etwas steigt, dann wird auch die Steuer-spannung am Verstärkertransistor größer. Sein

Kollektorstrom steigt, wodurch das Basispotential am Stelltransistor sinkt und sich dessen Durchlaßwiderstand erhöht. Damit wird einer Erhöhung der Ausgangsspannung entgegengewirkt. Sie bleibt nahezu konstant. Es liegt hier eine Regelung vor, bei der aus der Abweichung

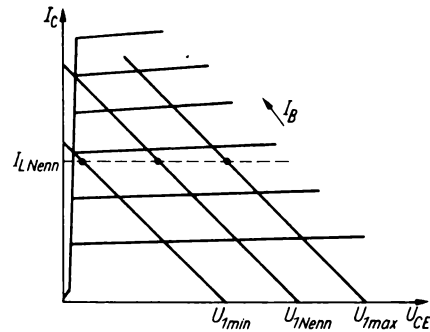


Bild 1.32. Arbeitspunkte des Stelltransistors bei konstanter Last und veränderlicher Eingangsspannung

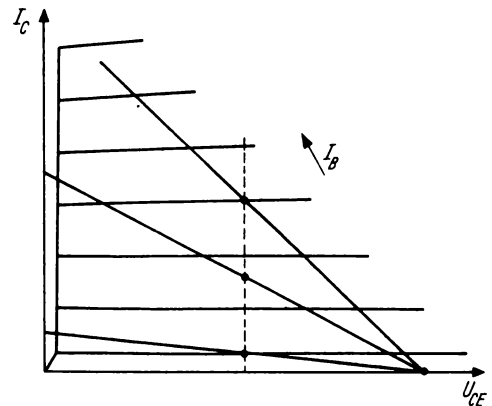


Bild 1.33. Arbeitspunkte des Stelltransistors bei konstanter Eingangsspannung und veränderlichem Lastwiderstand

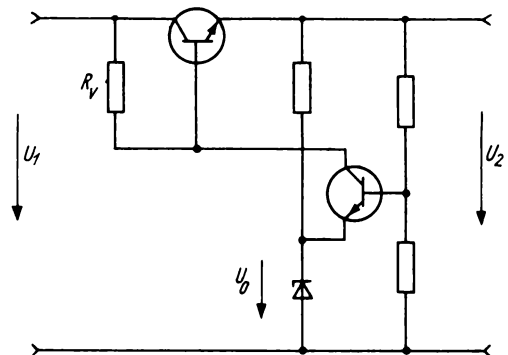


Bild 1.34. Grundsaltung des Spannungsstabilisators

der Ausgangsspannung vom Sollwert die für den Stelltransistor erforderliche Stellgröße gewonnen wird.

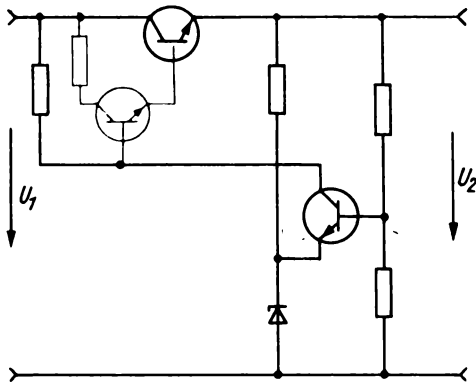


Bild 1.35. Spannungsstabilisator mit verbesserter Stabilisierung

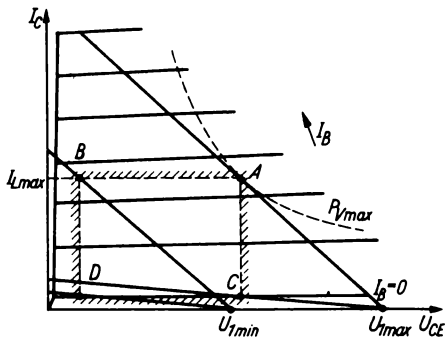


Bild 1.36. Vereinfachte Darstellung des Arbeitsbereichs des Stelltransistors

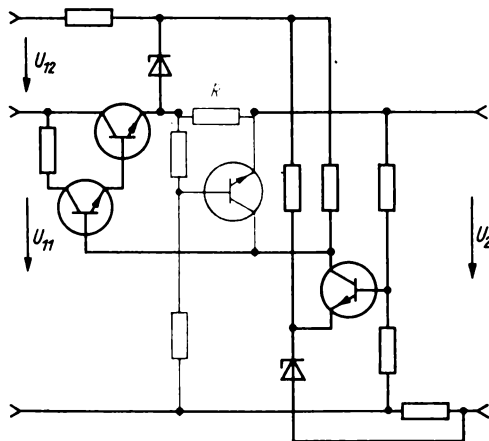


Bild 1.37. Stabilisierungsschaltung mit elektronischer Sicherung

Die Güte der Stabilisierung hängt wesentlich von der Verstärkung der Ausgangsspannungsdifferenz und von der Konstanz der Vergleichsspannung ab. Eine hohe Spannungsverstärkung des Verstärkertransistors setzt einen hochohmigen Arbeitswiderstand desselben voraus. Zur Anpassung dieses hochohmigen Widerstands an den niederohmigen Eingangswiderstand des Stelltransistors fügt man einen Transistor (oder mehrere) in die Schaltung ein (Bild 1.35).

Besondere Aufmerksamkeit ist der Temperaturabhängigkeit der Vergleichsspannung zu widmen. Niedrige Temperaturbeiwerte und differentielle Widerstände haben Z-Dioden mit einer Durchbruchspannung um 6 V. Des weiteren ist darauf zu achten, daß die Vergleichsspannung nur geringfügig niedriger als die stabilisierte Ausgangsspannung ist, damit möglichst die volle Ausgangsspannungsschwankung an der Basis des Verstärkertransistors wirksam wird.

Die am Stelltransistor auftretende Verlustleistung bestimmt den maximal möglichen Laststrom. Bild 1.36 zeigt den zu erwartenden Aussteuerbereich des Stelltransistors, abgeleitet aus der Vereinigung der Bilder 1.32 und 1.33. Die Punkte A und B ergeben sich für den größten Laststrom, die Punkte C und D für den geringsten Laststrom bei jeweils maximaler Eingangsspannungsschwankung. Der geringste Laststrom muß größer als der bei der größten Kristalltemperatur vorliegende Reststrom des Stelltransistors sein. Um auch bei Leerlauf eine einwandfreie Funktion der Schaltung zu haben, wird an die Ausgangsklemmen eine Vorlast angeschlossen.

Zum Schutz des Stelltransistors vor Überlastung wird eine elektronische Sicherung in die Schaltung eingefügt. Aus der Vielzahl der Schaltungsvarianten seien einige typische Vertreter ausgewählt. Bild 1.37 zeigt eine Schaltung, bei der ein zusätzlich in den Verstärker eingefügter Transistor den Überlastungsschutz übernimmt. Die Widerstände sind so eingestellt, daß bei normaler Belastung der Sicherungstransistor gesperrt ist. Steigt der Laststrom und damit der Spannungsabfall am Widerstand R über die zulässige Grenze an, so fließt durch den Sicherungstransistor ein Strom, wodurch der Basisstrom des Stelltransistors erniedrigt und dessen Durchlaßwiderstand erhöht wird. Ein Sinken der Ausgangsspannung ist die Folge. Bei ausgangsseitigem Kurzschluß ist zur Öffnung des Sicherungstransistors nur ein ge-

ringer Laststrom erforderlich, weil dann der volle Spannungsabfall am Widerstand als Steuerspannung wirksam wird. Es stellt sich ein wesentlich unter dem Nennlaststrom liegender Kurzschlußstrom ein. Die entstehende Lastkennlinie wird im Bild 1.38 gezeigt. Nach Wegfall der Überlast geht die Schaltung selbsttätig auf ihre ursprüngliche Ausgangsspannung zurück. Diese Kennlinie bezeichnet man als *Fold-back-Kennlinie*.

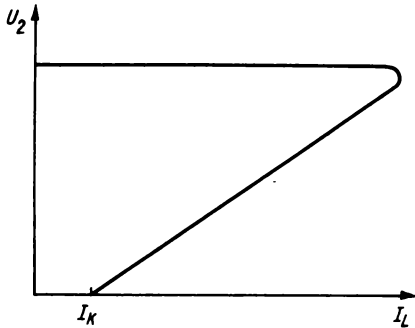


Bild 1.38. Lastkennlinie der Schaltung nach Bild 1.37

Eine weitere Schaltung arbeitet mit einem Flip-Flop (s. Abschn. 5.4.), der nach Überschreiten eines bestimmten Laststroms in seine zweite stabile Lage kippt und über einen zusätzlichen Entkopplungstransistor den Stelltransistor sperrt. Nach Wegfall der Last muß durch einen von Hand betätigten Taster der Flip-Flop zurückgeschaltet werden. Diese Schaltung gestattet die sichere Kontrolle, ob während des Betriebes eine Überlast vorgelegen hat.

Bei einer dritten Variante wirkt die Schaltung nach Erreichen eines Höchststroms stromstabilisierend. Mit kleiner werdendem Lastwiderstand geht bei konstant bleibendem Strom die Ausgangsspannung gegen 0. Bei dieser Schaltung steigt jedoch bei Kurzschluß die Verlustleistung am Stelltransistor erheblich an.

Aus der Darstellung im Bild 1.36 wird ersichtlich, daß bei hohen Lastströmen eine hohe Verlustleistung am Stelltransistor auftritt. Wendet man die Parallelregelung nach Bild 1.31 an, so entsteht gerade bei maximaler Last die minimale Verlustleistung am Stelltransistor. Diese Geräte sind kurzschlußfest. Die Wirkungsweise soll anhand des Bildes 1.39 beschrieben werden. Man geht von der Annahme aus, daß sich die Ausgangsspannung etwas erhöht. Der dann größer werdende Spannungsabfall am Vorwiderstand bewirkt ein Ansteigen des

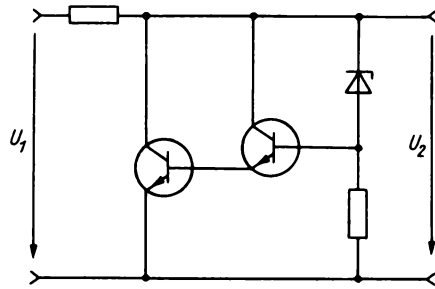


Bild 1.39. Spannungsstabilisierung mit Paralleltransistor

Basisstroms durch den Verstärkertransistor, wodurch auch der Kollektorstrom des Stelltransistors steigt. Der den Vorwiderstand durchfließende Strom wird deshalb größer, und der Spannungsabfall am Vorwiderstand steigt. Die Ausgangsspannung bleibt weitgehend konstant. Die maximale Verlustleistung am Paralleltransistor tritt bei ausgangsseitigem Leerlauf und maximaler Eingangsspannung auf.

Gewinnt man die Steuerspannung für den Verstärkertransistor nicht aus der Differenz der Ausgangsspannung mit der Vergleichsspannung, sondern aus der Differenz einer vom Laststrom abhängigen Spannung mit einer Vergleichsspannung, dann entsteht eine *Stromstabilisierung* (Stromkonstanthaltung). Das entspricht einer Stromquelle, die auch als Spannungsquelle mit sehr hohem Innenwiderstand aufgefaßt werden kann. Die Wirkungsweise soll anhand des Bildes 1.40 erläutert werden. Steigt der Laststrom, so steigt auch der Spannungsabfall am Widerstand R . Die Steuerspannung am Verstärkertransistor wird größer und bewirkt in bekannter Weise ein Sinken des Basisstroms am Stelltransistor. Der Kollektorstrom des Stell-

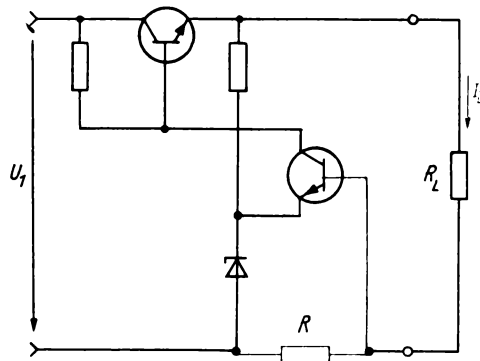


Bild 1.40. Stromstabilisierungsschaltung

transistors zeigt somit das Bestreben zu sinken. Der Laststrom bleibt weitgehend konstant. Eine elektronische Spannungsstabilisierung läßt sich auch problemlos mit integrierten Schaltkreisen realisieren. Bild 1.41 zeigt eine stark vereinfachte Darstellung. Besondere Aufmerksamkeit wird bei diesen integrierten Schaltkreisen der Erzeugung der Referenzspannung gewidmet. Die zugehörige Schaltung baut auf den Grundelementen der integrierten Verstärkerschaltungen auf, die erst im Abschn. 3. behandelt und hier nicht näher betrachtet werden. Die Referenzspannung wird dem nichtinvertierenden Eingang (mit positiver werdender Eingangsspannung wird auch die Ausgangsspannung positiver) eines Differenzverstärkers zugeführt. Über einen externen Spannungsteiler $R1/R2$ wird die geteilte Ausgangsspannung an den invertierenden Eingang (mit positiver werdender Eingangsspannung wird die Ausgangsspannung negativer) angeschlossen. Die verstärkte Differenz zwischen Ausgangs- und Referenzspannung wird zur Steuerung des Stelltransistors genutzt. Steigt beispielsweise die Ausgangsspannung, so steigt auch der Spannungsabfall an $R2$; der Ausgang des Differenzverstärkers wird negativer; der Durchlaßwiderstand des Stelltransistors steigt – damit wird einer Erhöhung der Ausgangsspannung entgegengewirkt. Die Schaltung unterscheidet sich in ihrem grundsätzlichen Aufbau nicht von der Schaltung nach Bild 1.34.

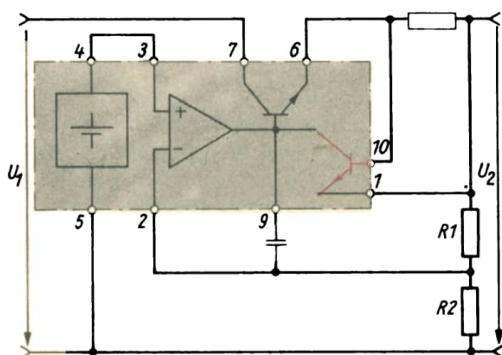


Bild 1.41. Schaltung zur elektronischen Spannungsstabilisierung mit integriertem Schaltkreis

Der im Bild 1.41 rot eingetragene, extern zugeschaltete Widerstand dient in Verbindung mit dem Transistor als Überlastungsschutz. Beim Erreichen eines bestimmten Ausgangsstroms wird der Strombegrenzungstransistor geöffnet.

Der ihn durchfließende Strom führt zu einer Erniedrigung des Basisstroms des Stelltransistors; die Schaltung wirkt jetzt stromstabilisierend. Der extern zuzuschaltende Kondensator dient der Frequenzgangkompensation (s. Abschnitt 3.4.). Für den Schaltkreis MAA 723 sind im Bild 1.41 die Schaltkreisanschlüsse eingetragen. Ist der zulässige Laststrom des Schaltkreises niedriger als der geforderte Laststrom, kann ein zusätzlicher Stelltransistor extern zugeschaltet werden (Bild 1.42).

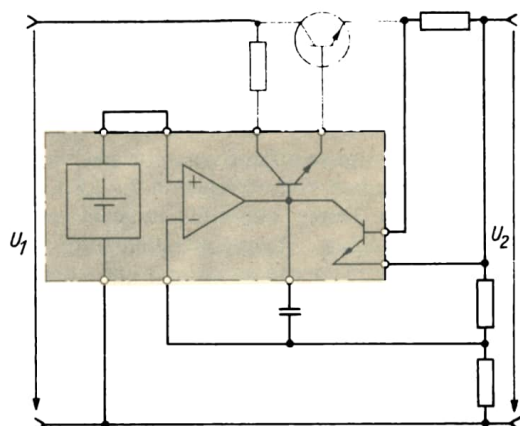


Bild 1.42. Zusätzlicher externer Stelltransistor zum Stabilisierungsschaltkreis

Die innere Stromversorgung des integrierten Schaltkreises kann an die zu stabilisierende Eingangsspannung angeschlossen werden, wenn diese in den für den integrierten Schaltkreis zulässigen Grenzen liegt. Ist dies nicht der Fall, so ist die geforderte Betriebsspannung gesondert zuzuführen. Ist eine stabilisierte Ausgangsspannung gefordert, die kleiner als die im integrierten Schaltkreis erzeugte Referenzspannung ist, so wird die Referenzspannung über einen Spannungsteiler geteilt dem Differenzverstärker zugeführt.

Eine für den Anwender besonders bequeme Variante stellen die integrierten Festspannungsregler und die Floatingregler dar. Sie verfügen über lediglich drei Anschlüsse (Eingang, Ausgang, Masse bzw. Einstellanschluß für die Ausgangsspannung) und benötigen nur eine minimale Außenbeschaltung. Gleichzeitig können verschiedene Schutzschaltungen integriert sein (z. B. Überstromschutz, thermischer Überlastschutz, Schutz vor Überspannungen). Eine Unterbringung auf der zu versorgenden Leiter-

platte wird damit bequem möglich. Im zentralen Stromversorgungssteil ist dann nur noch die Gleichrichterschaltung und der Ladekondensator vorzusehen.

Zu den Festspannungsreglern gehören z. B. die Schaltkreise MA 7805 und MA 7812 von Tesla (für 5 und 12 V). Die genannten Schaltkreise verfügen über einen internen Kurzschlußschutz mit einer Fold-back-Kennlinie und einen internen Wärmeschutz, der bei zu stark ansteigender Kristalltemperatur t_j (z. B. als Folge unzureichender Kühlung oder zu hoher Umgebungstemperatur t_a) die entnehmbare Leistung begrenzt. Bild 1.43 zeigt eine Schaltung mit einem derartigen Schaltkreis und der oft zweckmäßigen Außenbeschaltung. Der Parallelkondensator am Ausgang verbessert das Regelverhalten bei Lastschwankungen, sollte aber nach Herstellerangaben einen Wert von 2,2 μF nicht überschreiten. Treten, durch die zu versorgende Schaltung bedingt, höhere Kapazitätswerte zwischen Ausgangsklemme und Masse auf, ist zum Schutz des Schaltkreises die eingezeichnete Diode zwischen Ausgang und Eingang zu schalten. Der Kondensator zwischen Eingang und Masse (etwa 1 μF) vermindert eine eventuelle Schwingneigung des Schaltkreises. Beide Kondensatoren sind unmittelbar am Schaltkreis anzuordnen.

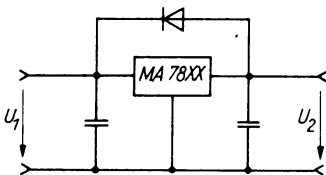


Bild 1.43. Spannungsstabilisierung mit Festspannungsregler

Der entnehmbare Ausgangsstrom beträgt im Beispiel nach Bild 1.43 1 A, wenn dabei die zulässige Verlustleistung noch nicht überschritten ist. In den meisten Fällen wird jedoch der entnehmbare Strom durch die zulässige Verlustleistung begrenzt. Höhere Ausgangsströme werden durch extern zugeschaltete Leistungstransistoren möglich. Bild 1.44 zeigt eine mögliche Variante. Erst wenn der durch den Schaltkreis fließende Strom einen vorgegebenen Wert erreicht hat, wird durch den Spannungsabfall an R_1 der Transistor V_1 geöffnet und übernimmt den zusätzlichen Strom. Um V_1 vor Kurzschluß zu schützen, ist der Widerstand R_2 eingefügt, dessen Spannungsabfall V_2 bei Stromüberschrei-

tung öffnet und damit den Spannungsabfall an R_1 wieder verringert (Stromerniedrigung durch V_1). Als Folge wirkt jetzt auch der interne Überlastschutz des Schaltkreises.

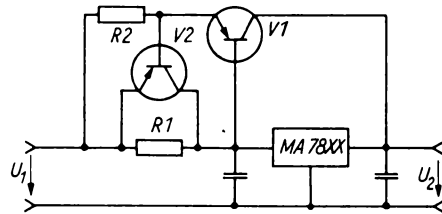


Bild 1.44. Festspannungsregler mit parallelgeschaltetem Leistungstransistor

Zu den Floatingreglern (engl., schwimmende Regler) gehören z. B. die Schaltkreise B 3170 und B 3370 vom Halbleiterwerk Frankfurt/O. Bei Ersterem handelt es sich um einen Positivregler (zur Regelung einer positiven Spannung), bei Zweiterem um einen Negativregler (zur Regelung einer negativen Spannung). Zwar unterscheiden sich die Innenschaltungen im Detail, das Grundprinzip und damit die erforderliche Außenbeschaltung ist jedoch weitgehend gleich. Bild 1.45 zeigt die Grundschaltung zur Spannungsstabilisierung und Bild 1.46 die innere Struktur des Schaltkreises. Im Inneren des Schaltkreises wird eine Referenzspannung von etwa 1,25 V erzeugt (Bandgap-Referenzquelle). Der invertierende Eingang des Differenzverstärkers ist mit dem Schaltkreisausgang verbunden. Am nichtinvertierenden Eingang wirkt bezogen auf den Schaltkreisausgang die Differenz aus Referenzspannung und Spannungsabfall am externen Widerstand R_1 . Da am hochverstärkten Differenzverstärker die Eingangsdifferenzspannung praktisch 0 ist, wird der Verstärker die Längstransistoren stets so steuern, daß der Spannungsabfall an R_1 gleich der Referenzspannung wird. So wird eine Erhöhung der Ausgangsspannung auch zu einer Erhöhung des

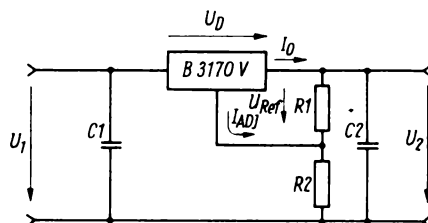


Bild 1.45. Spannungsstabilisierung mit Floatingregler

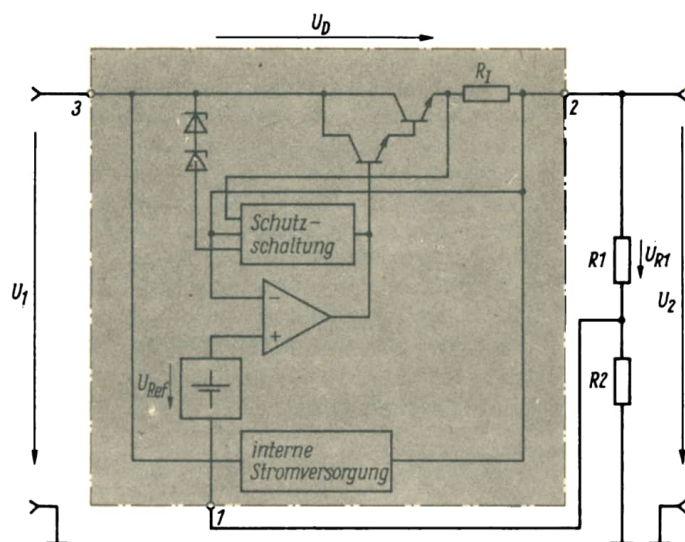


Bild 1.46. Innere Struktur des Floatingreglers B 3170

Spannungsabfall an $R1$ führen, und damit wird der nichtinvertierende Eingang des Differenzverstärkers negativer als der invertierende Eingang. Als Folge erniedrigt sich die Basis-Emitter-Spannung der Längstransistoren und die Ausgangsspannung sinkt.

Durch Teilung der Ausgangsspannung U_2 sind damit beliebige, über der Referenzspannung liegende Spannungswerte stabilisierbar. Der Spannungsabfall an $R1$ ist dabei immer gleich der Referenzspannung. Bei Vernachlässigung des Stromes I_{ADJ} folgt aus der Berechnung des Spannungsteilers $R1/R2$

$$U_2 = U_{Ref} (1 + R1/R2) = 1,25 \text{ V} (1 + R2/R1).$$

Wird dabei $R1$ mit 120Ω gewählt (Herstellerangabe), so wird der Ausgang mit dem funktionsbedingten Mindeststrom von 10 mA belastet und ist die Vernachlässigung von I_{ADJ} ($\leq 100 \mu\text{A}$) gerechtfertigt.

Wie auch beim Festspannungsregler sind unter Umständen weitere Außenbeschaltungen notwendig. Eine Verbesserung der Brummspannungsunterdrückung ist durch Zufügen von $C3$ (Bild 1.47) möglich. Wenn $C3$ oder $C2$ Kapazitätswerte größer als $22 \mu\text{F}$ haben, ist bei höheren Spannungen die Zuschaltung der Dioden $V1$ und $V2$ notwendig, damit der Schaltkreis durch die Entladeströme der Kondensatoren $C2$ und $C3$ beim Kurzschluß am Ein- oder Ausgang nicht zerstört wird.

Die Stromversorgung der internen Schaltungsteile wird aus der Differenzspannung zwischen U_1 und U_2 gewonnen (Bild 1.46). Daraus resul-

tiert auch die Forderung nach einer Mindestdifferenzspannung (im Beispiel 3 V).

Der Schaltkreis kann auch als Konstantstromquelle geschaltet werden. Bild 1.48 zeigt die entsprechende Grundschialtung. Zwischen Ausgang und Einstellanschluß liegt funktionsbedingt immer die Referenzspannung. Folglich muß der den Widerstand R durchströmende Strom stets so groß sein, daß an R die Referenzspannung abfällt. Unter Vernachlässigung des Stromes I_{ADJ} folgt

$$I_o = \frac{U_{Ref}}{R} = k$$

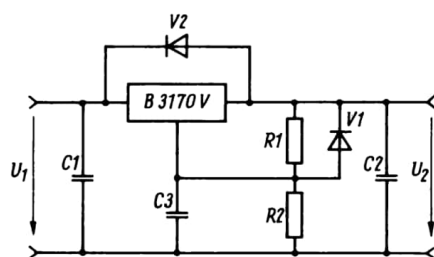


Bild 1.47. Komplette Stabilisierungsschaltung mit B 3170

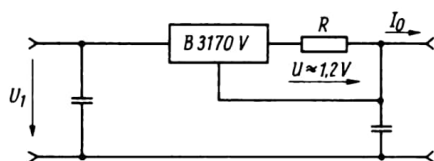


Bild 1.48. Schaltung zur Stromstabilisierung mittels Floatingreglers

(Im Beispiel ist I_0 in einem Bereich von 10 mA bis 1,5 A mit Widerstandswerten von 120 Ω bis 0,85 Ω realisierbar).

Die interne Schutzschaltung (Bild 1.46) vereinigt den Überstromschutz (durch Überwachung des stromabhängigen Spannungsabfalles am internen Widerstand R_I), Schutz gegen zu hohe Differenzspannung (angesteuert über 2 Z-Dioden) sowie den thermischen Überlastschutz und arbeitet auf den Längstransistor.

Schaltregler

Die bisher betrachteten *stetigen Regler* haben den Vorteil, daß sie auch äußerst kurzzeitige Laständerungen ausregeln können (geringe Regelzeitkonstante) und eine äußerst geringe Restwelligkeit der geregelten Gleichspannung aufweisen. Dem steht die hohe Verlustleistung am Längstransistor als entscheidender Nachteil gegenüber. Eine um etwa zwei Drittel niedrigere Verlustleistung ermöglichen *Schaltregler*.

Bild 1.49 zeigt das Prinzip eines Schaltreglers. Der vom stetigen Regler bekannte Längstransistor wird hierbei durch einen als Schalter wirkenden Transistor ersetzt. Die Gleichspannung U_1 wird rhythmisch geschaltet. Dabei kann die Schaltfrequenz konstant sein und die Einschaltdauer verändert werden (auch als *Pulsdauermodulation* PDM bezeichnet) oder die Einschaltdauer konstant sein und die Schaltfrequenz verändert werden (auch als *Pulsfrequenzmodulation* PFM bezeichnet). Wegen der einfacheren Entstörung wird die PDM bevorzugt eingesetzt. Entsprechende Ansteuerschaltkreise stehen heute als integrierte Schaltkreise zur Verfügung (z. B. der integrierte Schaltkreis B 260 vom VEB Halbleiterwerk Frankfurt/Oder). Die Schaltfrequenz wird etwa 16...50 kHz gewählt; beim Einsatz von V-MOSFET sind Schaltfrequenzen bis 500 kHz möglich.

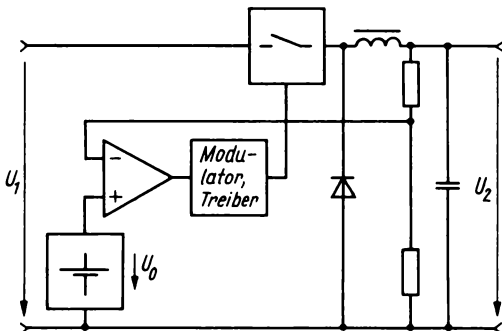


Bild 1.49. Prinzip des Schaltreglers

Die Verluste am Schalttransistor setzen sich zusammen aus den Verlusten im Ein-Zustand ($U_{CEsat} I_C$), den verhältnismäßig kleinen Verlusten im Aus-Zustand und den Umschaltverlusten. Besonders letztere begrenzen die Schaltfrequenz nach oben.

Zur Erklärung der Wirkungsweise der Schaltungen wird zunächst wiederholend das Schaltverhalten von Spulen betrachtet. Bild 1.50 a zeigt die bekannte Schaltung und Bild 1.50 b den Strom- und Spannungsverlauf beim Schalten einer Spule an Gleichspannung. In dem Diagramm ist die Zeitkonstante angegeben, die bei *RL-Schaltungen* nach

$$\tau = \frac{L}{R} \quad (1.16)$$

berechnet wird. Man erkennt, daß es an Spulen keine sprunghaften Stromänderungen gibt.

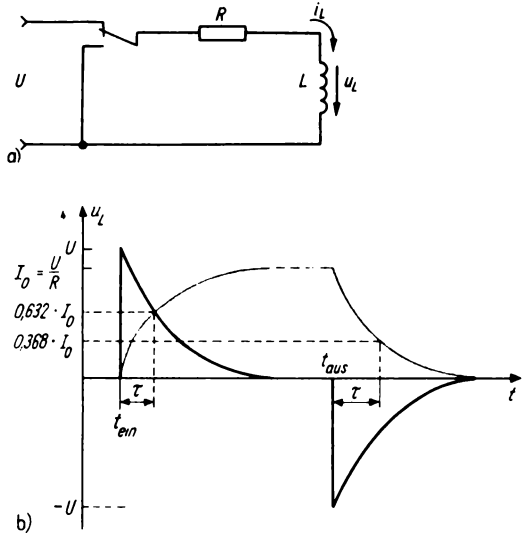


Bild 1.50. Zu- und Abschalten einer Spule an Gleichspannung

a) Schaltung; b) Strom- und Spannungsverlauf an der Induktivität

Führt man das Zu- und Abschalten der *RL-Schaltung* in einer gegenüber τ sehr kurzen Zeit durch, so wird jeweils nur ein kleiner Teil der exponentiellen Kurven durchlaufen. Diesen Teil kann man mit hinreichender Genauigkeit als linear betrachten. Eine ausführlichere Darstellung der analogen Verhältnisse bei *RC-Schaltungen* ist im Abschn. 4.3. gegeben. Schaltet der Schalttransistor (s. Bild 1.49) ein,

so fließt durch die Spule ein ansteigender Strom. Dabei wird im sich aufbauenden Magnetfeld Energie gespeichert. Gleichzeitig wird auch der Kondensator aufgeladen. Beim Abschalten des Schalttransistors wird die in der Spule gespeicherte magnetische Energie wieder in elektrische Energie gewandelt und über die Diode dem Verbraucher zugeführt (s. Abschn. 2.; Bild 2.5). Die Induktivität der Spule muß dem Laststrom angepaßt werden. Sollte die Belastung wegfallen, so wird die Spule wirkungslos (τ ginge dann gegen 0), und am Ausgang läge im Schaltmoment die volle Eingangsspannung. (Als Folge der Eigeninduktivität des Kondensators könnte diese Spannungsspitze am Ausgang nicht hinreichend schnell vom Kondensator begrenzt werden.) Die Schaltungen sind deshalb vor ausgangsseitigem Leerlauf zu schützen. Die Eingangsspannung wird zweckmäßig 3- bis 6mal höher als die Ausgangsspannung gewählt. Bei veränderter Anordnung von Schalttransistor und Spule sind allerdings auch Werte der Ausgangsspannung, die über den Werten der Eingangsspannung liegen, möglich.

Die hohe Schaltfrequenz erfordert den Einsatz von schnellen Dioden, die zur Vermeidung von HF-Störungen soft-recovery-Verhalten haben sollten.

Sollten die Nachteile des Schaltreglers (hohe Regelzeitkonstante, verhältnismäßig hohe Restwelligkeit der Ausgangsspannung) für einen bestimmten Anwendungsfall nicht vertretbar sein, kann man dem Schaltregler einen stetigen Regler mit geringem Spannungsabfall am Längstransistor nachschalten.

Sowohl beim stetigen Regler als auch beim Schaltregler führt ein Kurzschluß im Längstransistor zu einer unzulässig hohen Ausgangsspannung. Werden durch die Stromversorgung empfindliche Schaltkreise gespeist, so könnte das eine Zerstörung aller gespeisten Schaltkreise bewirken. (Für die im Abschn. 5. beschriebenen TTL-Schaltkreise gilt als absoluter Grenzwert für die mit 5 V angegebene Betriebsspannung 7 V.) Zum Schutz vor derartigen Überspannungen schaltet man parallel zu den Ausgangsklemmen des Reglers einen Thyristor, der bei Spannungsüberschreitung geöffnet wird und so die Stromversorgung kurzschließt. Als Folge unterbricht dann die vorgeschaltete Schmelzsicherung den Stromkreis. Die jeweils auf den Längstransistor wirkende elektronische Sicherung ist bei dessen Defekt wirkungslos.

In Tafel 1.2 sind die Stabilisierungsschaltungen noch einmal gegenübergestellt.

Zusammenfassung

Bauelemente mit einem fast senkrecht verlaufenden Teil der Strom-Spannungs-Kennlinie eignen sich gut zur Spannungsstabilisierung. Die stabilisierende Wirkung tritt in Verbindung mit einem Vorwiderstand durch die starke Stromzunahme bei geringer Spannungsänderung auf. Bei der Berechnung des Vorwiderstandes ist unter den jeweils ungünstigsten Bedingungen zu beachten, daß der maximale Querstrom nicht überschritten und der minimale Querstrom nicht unterschritten wird. Die vorgeschriebenen Kühlbedingungen sind zu beachten. Zur Erzielung guter Stabilisierungseigenschaften wird stets der größtmögliche Vorwiderstand eingesetzt.

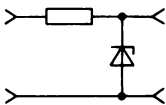
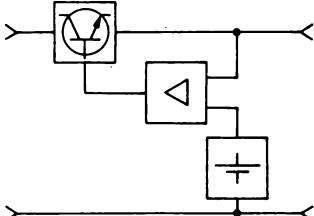
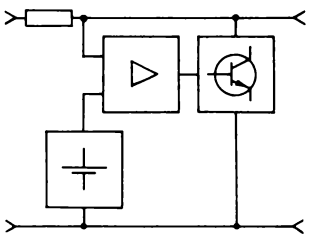
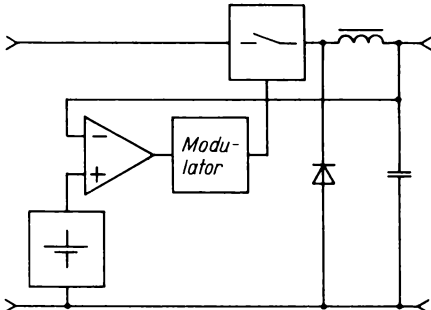
Soll eine berechnete Schaltung aufgebaut werden, so ist der zulässige Toleranzbereich für den Vorwiderstand zu bestimmen und ein Widerstand auszuwählen, dessen Toleranzfeld innerhalb des berechneten Toleranzfeldes liegt.

Mit Schaltungen zur elektronischen Spannungsstabilisierung sind Netzgeräte herstellbar, deren Innenwiderstand gegen Null geht. Prinzipiell wird die Ausgangsspannung oder ein Teil derselben mit einer Referenzspannung verglichen und die Differenz zur Steuerung des Stelltransistors (meist über einen mehrstufigen Gleichspannungsverstärker) benutzt.

Am Stelltransistor tritt eine erhebliche Verlustleistung auf. Die elektronische Sicherung schützt ihn vor Überlastung. Entweder schaltet diese Sicherung die Ausgangsspannung ab (erneute Einschaltung nach Beseitigung der Überlast von Hand erforderlich), oder sie begrenzt den Strom auf einen für den Stelltransistor ungefährlichen Wert bzw. sorgt nach Überschreiten des Höchststroms für Stromstabilisierung. Schaltungen mit Paralleltransistor sind kurzschlußfest, aber nur für kleine Ausgangsspannungen wirtschaftlich. Eine Stromstabilisierung entsteht, wenn eine dem Laststrom proportionale Spannung mit einer Vergleichsspannung verglichen wird und die hierbei entstehende Differenz den Stelltransistor steuert.

Immer häufiger werden wegen ihrer bequemen Einsetzbarkeit integrierte Festspannungsregler oder Floatingregler verwendet. Schaltregler weisen einen wesentlich besseren Wirkungsgrad als stetige Regler auf. Ein als Schalter betriebener Längstransistor wird entweder durch pulsdauer-

Tafel 1.2. Stabilisierungsschaltungen

Bezeichnung	Schaltung	Eigenschaften
Stabilisierungsschaltung mit Z-Diode		einfacher Aufbau, Einsatz für niedrige Spannungen bei geringen Anforderungen
Stabilisierungsschaltung mit Reihenregelung		äußerst hohe Stabilisierungswirkung erreichbar, empfindlich gegen Überlast, elektronische Sicherung erforderlich, Einsatz für niedrige und mittlere Spannungen
Stabilisierungsschaltung mit Parallelregelung		gute Stabilisierungswirkung, unempfindlich gegen Überlast, Einsatz für niedrige Spannungen
Schaltregler		hoher Wirkungsgrad, hohe Regelzeitkonstante, verhältnismäßig hohe Restwechselspannung

oder durch pulsfrequenzmodulierte Impulse geschaltet, die in einem speziellen Modulator aus dem Vergleich der Ausgangsspannung mit der Referenzspannung gewonnen werden. Nachteilig ist die große Regelzeitkonstante und die höhere Restwelligkeit der Ausgangsspannung.

Zum Schutz vor ausgangsseitiger Überspannung beim Kurzschluß des Längstransistors kann den Ausgangsklemmen ein Thyristor parallel geschaltet werden, der bei Überspannung die Stromversorgung kurzschließt.

1.5. Transverter

1.5.1. Allgemeines

Eine elektronische Schaltung, mit deren Hilfe man eine Gleichspannung in eine andere wandeln kann, nennt man *Transverter* oder *Gleichspannungswandler* (auch DC-DC-Wandler entsprechend der englischen Bezeichnung für Gleichstrom: direct current). Ihr Einsatz erfolgt überall dort, wo in batteriegespeisten Geräten eine gegenüber der Batteriespannung höhere

Betriebsgleichspannung benötigt wird (z. B. in transportablen Sendeanlagen, in batteriegespeisten Blitzgeräten) oder im Schaltnetzteil, um aus einer hohen Gleichspannung eine niedrigere Gleichspannung zu gewinnen.

Eine Änderung der Spannung setzt deren Transformation voraus. Da nur eine Wechselspannung oder eine einer Gleichspannung überlagerte Wechselspannung transformierbar ist, muß im Transverter eine *Wechselrichtung* erfolgen. Würde man hierzu die im Abschn. 4.2. beschriebenen Sinusoszillatoren verwenden, so wäre der maximal erreichbare Wirkungsgrad nur 40 %. Einen Wirkungsgrad bis zu 80 % kann man erreichen, wenn man die Gleichspannung in eine Rechteckspannung wandelt und diese transformiert. Der Transistor arbeitet dann als *elektronischer Schalter* (s. auch Abschn. 5.2.).

Nach dem Zeitpunkt der Stromentnahme auf der Sekundärseite unterscheidet man die Transverter in

- **Sperrwandler** Entnahme der Energie während der Sperrzeit der Primärseite
- **Summierwandler** Entnahme der Energie sowohl während der Sperrzeit als auch während der Stromflußzeit der Primärseite
- **Stromflußwandler** Entnahme der Energie während der Stromflußzeit der Primärseite.

Für höhere Leistungen wird der Stromflußwandler auch in einer Gegentaktschaltung betrieben.

1.5.2.

Sperrwandler

Die grundsätzliche Wirkungsweise des Sperrwandlers läßt sich anhand des Bildes 1.51. erklären. Wird der Schalter geschlossen, so wird

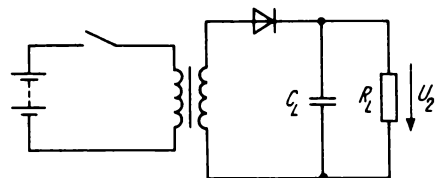


Bild 1.51. Prinzip eines Sperrwandlers

in der Primärspule ein ansteigender Strom fließen. Innerhalb der Spule baut sich ein Magnetfeld auf und speichert die Energie.

$$W = \frac{L}{2} I^2 \quad (1.17)$$

Gleichzeitig wird während des Stromanstiegs in der Sekundärwicklung eine Spannung induziert, für die der Gleichrichter in Sperrrichtung liegt. Öffnet man jetzt den Schalter, so wandelt sich die im Kern gespeicherte magnetische Energie wieder in elektrische Energie um. Für die dabei in der Sekundärwicklung induzierte Spannung liegt der Gleichrichter in Durchlaßrichtung, und der Kondensator wird aufgeladen. Die Funktion des Schalters kann ein Transistor übernehmen (Bild 1.52). Zu seiner Steuerung dient eine dritte Wicklung des Transformators. Wird bei dieser Schaltung die Batterie zugeschaltet, so fließt durch die Wicklung $L1$ ein linear ansteigender Strom. (Vorausgesetzt, daß die Stromflußzeit sehr klein gegenüber τ ist.) Dadurch wird in der Wicklung $L3$ nach Gl. (1.18)

$$U = \frac{\Delta I}{\Delta t} L \quad (1.18)$$

eine konstante Spannung induziert, die so gerichtet ist, daß sie den Transistor öffnet. Der nun fließende Basisstrom bestimmt den Höchstwert, auf den der Strom in der Wicklung $L1$ steigen kann. Ist dieser Höchstwert erreicht, so liegt keine weitere Stromänderung mehr vor. In der Wicklung $L3$ wird keine Spannung mehr induziert, und der Transistor ist gesperrt. Als Folge bricht das vorher aufgebaute Magnetfeld zusammen und induziert in der Wicklung $L2$ eine Spannung, für die der Gleichrichter in Durchlaßrichtung liegt. Ist diese Spannung größer als die Ladespannung am Kondensator, so wird er aufgeladen (vgl. Bild 1.15). Gleichzeitig wird auch in $L3$ eine Spannung induziert, die

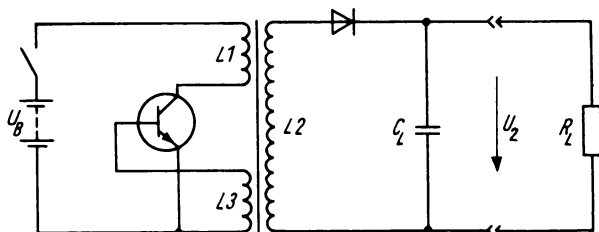


Bild 1.52. Sperrwandler

jetzt so gerichtet ist, daß sie den Transistor völlig sperrt. Da auch in $L1$ beim Zusammenbrechen des Magnetfelds eine Spannung induziert wird, erhöht sich die am Transistor liegende Spannung auf die Summe aus Batteriespannung und Induktionsspannung.

Die Arbeitsweise des Sperrwandlers ist im Bild 1.53 anhand des Kennlinienfeldes dargestellt. Im Einschaltmoment liegt der Arbeitspunkt A vor. Nach dem oben beschriebenen Vorgang folgt aus dem einsetzenden linearen Stromanstieg ein konstanter Basisstrom. Auf der dem Basisstrom I_B zugeordneten Kennlinie läuft jetzt der Arbeitspunkt über B mit steigendem Kollektorstrom nach oben bis zur Begrenzung des Kollektorstromanstiegs durch den vorgegebenen Basisstrom (Punkt C). Damit ist der Abschaltmoment erreicht. Da jetzt in der Wicklung $L3$ eine entgegengerichtete Spannung induziert wird, springt der Arbeitspunkt sehr schnell auf D und wandert mit der in $L1$ induzierten Spannung nach E. Da die Basis in Sperrrichtung vorgespannt ist, liegt der Arbeitspunkt unterhalb der Reststromkennlinie. Mit dem Abklingen des Abbaus des Magnetfelds kehrt der Arbeitspunkt nach A zurück. Der gesamte Vorgang beginnt von vorn. Da der Übergang vom Arbeitspunkt C nach D in sehr kurzer Zeit erfolgt, darf die Verbindungslinie beider Arbeitspunkte die Verlustleistungshyperbel schneiden. Die Darstellung verdeutlicht, daß der Transistor eine wesentlich über der Batteriespannung liegende Sperrspannung haben muß.

Bei der beschriebenen Schaltung wurde der Abschaltmoment durch die durch I_B gegebene Begrenzung des Kollektorstroms bestimmt. Eine andere Form der Begrenzung ist die *Sättigungsbegrenzung*. Hierbei wird bei Erreichen der Sättigung des Kernes keine der Kollektorstromzunahme proportionale Flußänderung mehr erfolgen. Die in $L3$ induzierte Spannung sinkt. Der Abschaltmoment ist erreicht. Die vereinfachten Strom- und Spannungsverläufe werden im Bild 1.54 gezeigt.

In der beschriebenen Schaltung ist nicht gewährleistet, daß der Transverter anschwingt. Im Einschaltmoment muß ein geringer Basisstrom fließen, damit ein erstes Ansteigen des Kollektorstroms möglich wird. Als Starthilfe kann man einen Basisspannungsteiler nach Bild 1.55 einsetzen. Dieser ist so einzustellen, daß ein gerade ausreichender Basisstrom für ein sicheres Anschwingen fließt. Da dieser Spannungsteiler den Wirkungsgrad verschlechtert, wird er häufig

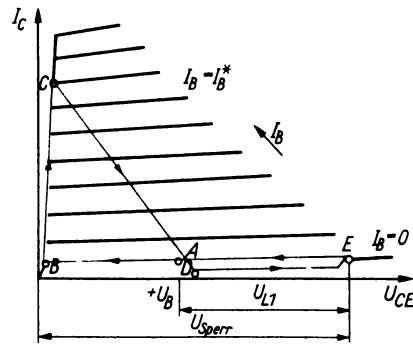


Bild 1.53. Arbeitsweise eines Sperrwandlers im Kennlinienfeld

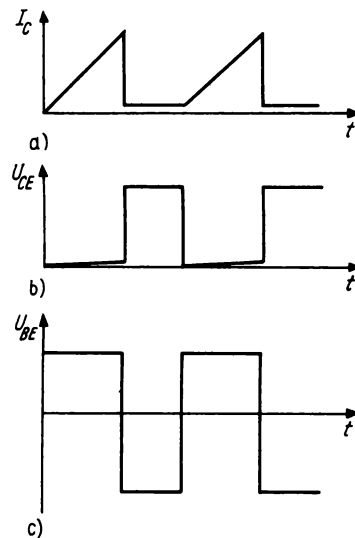


Bild 1.54. Strom- und Spannungsverläufe am Sperrwandler

a) Kollektorstromverlauf; b) Kollektorspannungsverlauf;
c) Basisspannungsverlauf

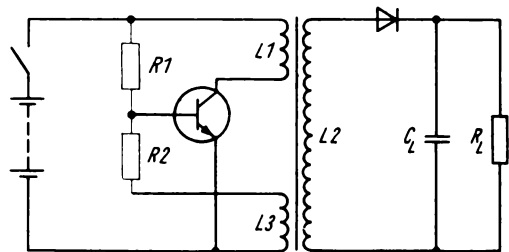


Bild 1.55. Sperrwandler mit Basisspannungsteiler zum sicheren Anschwingen

nur über eine Starttaste, die eventuell mit dem Einschalter verbunden ist, kurzzeitig eingeschaltet.

Aus der Funktionsbeschreibung geht hervor, daß die auf der Sekundärseite zur Verfügung stehende Energie nur von der je Schwingperiode im Kern gespeicherten Energie abhängt. Nach

$$W = \frac{U_2^2}{R_L} t \quad (1.19)$$

folgt für die Ausgangsspannung

$$U_2 = \sqrt{\frac{WR_L}{t}} \quad (1.20)$$

Dieser Gleichung ist zu entnehmen, daß die Ausgangsspannung stark lastabhängig ist. Bei sekundärseitigem Kurzschluß wird die Sekundärspannung Null. Die Primärseite wird dabei nicht überlastet. Die Sekundärseite ist mit einer Spannungsquelle mit sehr hohem Innenwiderstand vergleichbar.

Aus Gl. (1.20) erkennt man, daß bei Leerlauf eine sehr hohe Spannung auftreten wird. Die gesamte im Kern gespeicherte Energie wird dann an den Wicklungen $L1$ und $L3$ wirksam. Die dabei entstehenden sehr hohen Induktionsspannungen würden zur Zerstörung des Transistors führen. Ist kein sicherer Schutz gegen Lastausfall möglich, so kann man Überspannungsbegrenzer (Varistoren oder Glimmlampen) in die Schaltung einfügen. Eine bei Akkumulatorbetrieb besonders günstige Schaltung zur Leerlaufspannungsbegrenzung wird im Bild 1.56 gezeigt. Über eine weitere Wicklung wird ein Teil der Energie auf die Batterie zurückgeführt. Die in diesem Zweig liegende Diode ist durch die von der Batteriespannung abhängige Vorspannung so lange gesperrt, bis durch Ver-

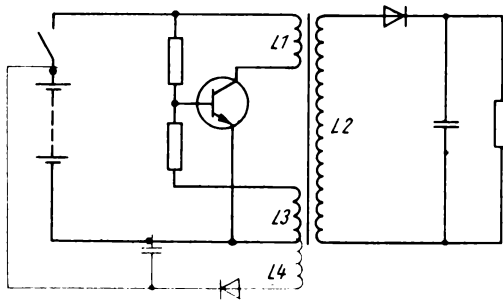


Bild 1.56. Sperrwandler mit Leerlaufspannungsbegrenzung

minderung der sekundärseitigen Last die Ausgangsspannung stark ansteigt.

Da die Sperrwandler, wie auch alle anderen Transverter, mit hoher Schaltfrequenz arbeiten (meist oberhalb des Hörbereichs), ist der Aufwand für die Glättung der pulsierenden Gleichspannung wesentlich geringer als bei der Netzfrequenz, vgl. Gln. (1.6), (1.8) und (1.9).

Eine Sonderform stellen geregelte Sperrwandler dar. Die Regelung erfolgt durch eine Beeinflussung des Endwertes des Stromes in der Stromflußphase und damit der je Periode im Kern gespeicherten Energie. Das Abschalten muß dabei durch eine in Abhängigkeit von der Sekundärspannung arbeitende Schaltung erzwungen werden. Um ein schnelles Sperren des Schalttransistors zu erreichen, kann man in die Schaltung einen „Ausräumkondensator“ (C auf Bild 1.57)

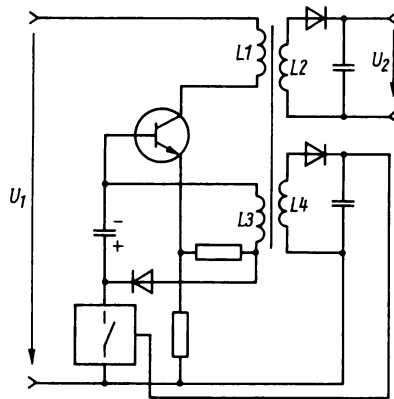


Bild 1.57. Prinzip des geregelten Sperrwandlers

einfügen. Während der Sperrphase wird dieser Kondensator in der eingetragenen Polarität aufgeladen. In der Flußphase ist er durch die Diode unwirksam und der Transistor wird durch $L3$ geöffnet. Soll jetzt der Stromfluß abgeschaltet werden, wird der Schalter geschlossen und damit die Basis des Schalttransistors durch den Kondensator negativ vorgespannt. Es kommt zu einem schnellen „Ausräumen“ der Basiszone. Die Abschaltung des Kollektorstromes bewirkt nun wieder die Induktion einer Spannung in $L3$, die den Transistor weiterhin sperrt. Der Schalter S wird geöffnet und C erneut aufgeladen. Eine erneute Stromflußphase schließt sich an.

1.5.3.

Summierwandler

Bild 1.58 zeigt eine Summierwandlerschaltung. Die Sekundärseite entspricht schaltungstechnisch einer Spannungsverdopplerschaltung nach Bild 1.10. Es wird auch die während der Stromflußphase der Primärseite in der Sekundärseite induzierte Spannung genutzt (Diode 2). Der Kollektorstrom setzt sich folglich aus der Summe des linear ansteigenden Stromes mit dem transformierten Laststrom durch Diode 2 zusammen. Der Abschaltmoment der Primärseite wird auch bei dieser Schaltung durch die Begrenzung des Kollektorstroms durch I_B^* festgelegt. Vorteilhaft ist die geringere Lastabhängigkeit der Spannung und der für gleiche Ausgangsleistung gegenüber der Sperrwandlerschaltung nur etwa halb so große Kollektorspitzenstrom.

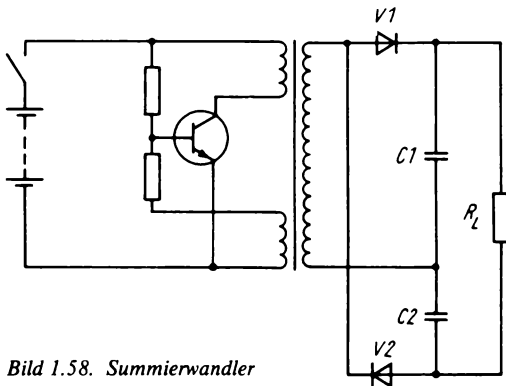


Bild 1.58. Summierwandler

Die Verdopplerschaltung wird dabei eingesetzt, weil die in der Stromflußphase und in der Sperrphase der Primärseite in der Sekundärwicklung induzierten Spannungen ungleich sind. Eine Zweiweggleichrichtung ist für solche Fälle nicht einsetzbar.

1.5.4.

Stromflußwandler

Die Schaltung des Stromflußwandlers (Bild 1.59) ist der Schaltung des Sperrwandlers ähnlich. Der Gleichrichter ist jetzt so gepolt, daß die während der Stromflußphase in der Wicklung L_2 induzierte Spannung ausgenutzt wird. Die Begrenzung des primären Stromanstiegs erfolgt auch hier durch I_B^* . Im Abschaltmoment würde eine sehr hohe Induktionsspannung entstehen (s. „Leerlauf des Sperrwandlers“); des-

halb schaltet man auf der Sekundärseite einen Kondensator C_0 parallel zur Wicklung L_2 . Dieser nimmt die im magnetischen Feld gespeicherte Energie auf. Es entsteht eine aperiodische Entladung. Die Ausgangsspannung des Stromflußwandlers wird wesentlich durch das Übersetzungsverhältnis des Transformators bestimmt. Die Ausgangsspannung ist weniger lastabhängig als bei den vorherigen Schaltungen.

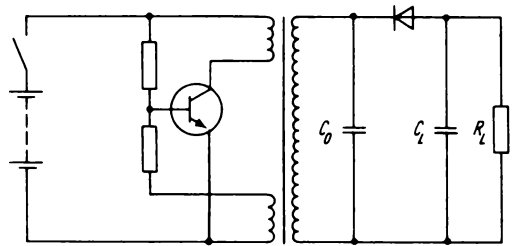


Bild 1.59. Stromflußwandler

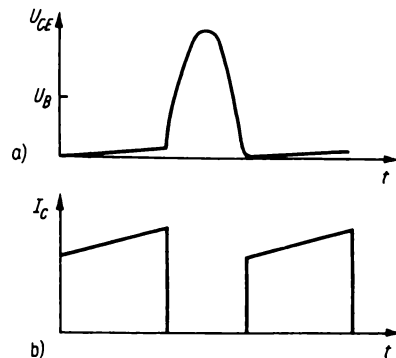


Bild 1.60. Strom- und Spannungsverlauf am Stromflußwandler

a) Kollektorspannungsverlauf; b) Kollektorstromverlauf

Der zeitliche Verlauf von I_C und U_{CE} ist im Bild 1.60 dargestellt.

Eine besondere Form des Stromflußwandlers ist der *Gegentaktwandler* (Bild 1.61). Er wird bei hohen Leistungen eingesetzt. Die Basiswicklungen sind so gepolt, daß der Transistor, durch den der linear ansteigende Strom fließt, die notwendige Basisspannung erhält. In der Basiswicklung für den Gegentransistor wird dann eine Spannung induziert, die diesen völlig sperrt. Auch hier ist der Abschaltmoment durch die Begrenzung des Kollektorstromanstiegs mit Hilfe des vorgegebenen Basisstroms bestimmt. Bricht jetzt das Magnetfeld zusammen, so wird der Gegentransistor leitend, und es baut sich ein Magnetfeld in um-

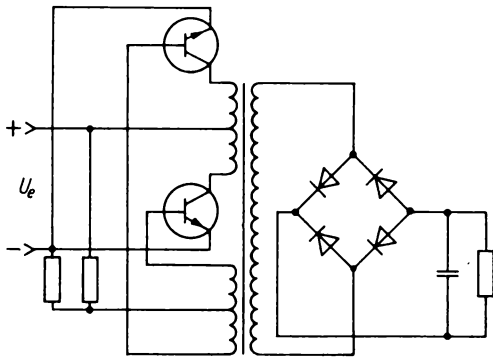


Bild 1.61. Gegentaktwandler

gekehrter Polarität auf. Es erfolgt wiederum Begrenzung und Umkehrung des Vorgangs.

Die auf der Sekundärseite induzierte Spannung wird in einer Zweiweggleichrichterschaltung gleichgerichtet.

Eine gleichmäßige Lastaufteilung auf beide Transistoren entsteht nur bei annähernd gleichen Daten der Transistoren.

Die Ausgangsspannung des Gegentaktwandlers ist weitgehend lastunabhängig. Bei Leerlauf tritt nur eine geringfügige Erhöhung der an den Bauelementen liegenden Sperrspannung ein. Die Schaltung ist mit einer Spannungsquelle mit niedrigem Innenwiderstand vergleichbar.

Zusammenfassung

Im Transverter wird eine Gleichspannung in eine Rechteckspannung gewandelt, transformiert und wieder gleichgerichtet. Der Transistor arbeitet dabei als elektronischer Schalter. Nach dem Zeitpunkt der Stromentnahme unterteilt man die Schaltungen in Sperrwandler, Summierwandler und Stromflußwandler. Die einzelnen Arten zeigen eine unterschiedliche Abhängigkeit der Ausgangsspannung von der Last. Der Sperrwandler ist mit einer Spannungsquelle mit hohem Innenwiderstand und der Stromflußwandler mit einer Spannungsquelle mit niedrigem Innenwiderstand vergleichbar. Der Sperrwandler ist kurzschlußfest und gegen Leerlauf, der Stromflußwandler leerlaufest und gegen Überlastung zu schützen. Bei allen Schaltungen erfolgt die Beendigung der Stromflußphase durch die Begrenzung des Kollektorstromanstiegs durch I_B . Für geforderte große Ausgangsleistungen wird der nach dem Stromflußwandlerprinzip arbeitende Gegentaktwandler eingesetzt.

1.6. Schaltnetzteile

Bei der konventionellen Stromversorgung mit dem Aufbau Transformator-Gleichrichter-Glättung des pulsierenden Gleichstroms (–Stabilisierung der Gleichspannung) wirkt der erforderliche Netztransformator volumen- und massebestimmend. Das *Schaltnetzteil* (auch mit einem Schaltregler kombinierbar) benötigt hingegen keinen Netztransformator und ermöglicht einen Wirkungsgrad bis $\eta = 0,9$. Rohstoff- und Energieeinsparungen werden möglich.

Das Prinzip eines Schaltnetzteils ist im Bild 1.62 dargestellt. Die Netzwechselspannung wird ohne vorherige Transformation gleichgerichtet und mit einem Ladekondensator geglättet. Die entstandene hohe Gleichspannung wird einem Transverter zugeführt. Durch geeignete Wahl der Windungszahlen des Übertragers im Transverter kann auf der Sekundärseite die gewünschte niedrige Gleichspannung entnommen werden. Die Schaltfrequenz des Transverters wird größer als 15 kHz gewählt. Deshalb kann ein verhältnismäßig kleiner Übertrager eine hohe Leistung übertragen (z. B. ein Übertrager mit dem Ferritkern EE 42 etwa 150 W).

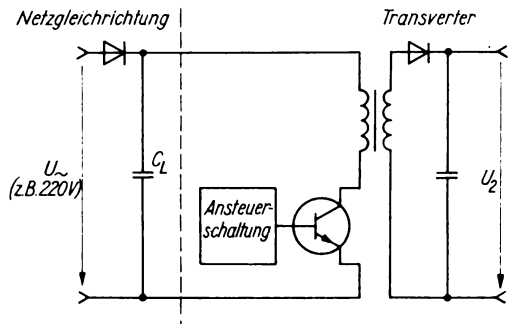


Bild 1.62. Prinzip eines Schaltnetzteils

Besonders hohen Anforderungen muß der Schalttransistor hinsichtlich seiner zulässigen U_{CE} gerecht werden. In diesem Zusammenhang wird vielfach zur genaueren Beschreibung der Grenzwerte des Schalttransistors das SOAR-Diagramm angegeben [safe operating area (engl.) sicheres Arbeitsgebiet].

Die Ansteuerung des Schalttransistors erfolgt heute meist fremdsynchronisiert über eine integrierte Ansteuerschaltung, die dem Schalttransistor – wie auch beim Schaltregler – Impulse mit konstanter Frequenz und unterschiedlicher

Impulsdauer (PDM) oder Impulse mit konstanter Impulsdauer und unterschiedlicher Frequenz (PFM) zuführt. Damit ist ein Schaltnetzteil mit gleichzeitiger Schaltreglerfunktion möglich.

Probleme treten bei der geforderten galvanischen Trennung zwischen Netz und erzeugter Gleichspannung auf, da bei gleichzeitiger Reglerfunktion die Spannung U_2 über eine Ansteuerschaltung mit dem netzseitig eingefügten Schalttransistor verbunden werden muß. Durch Zwischenschaltung eines Optokopplers oder eines Übertragers ist dieses Problem lösbar. Bild 1.63 zeigt das Prinzip einer entsprechenden Schaltung. Die erforderliche Betriebsspannung für die Ansteuerschaltung, den Differenzverstärker und die Referenzspannungsquelle wird zweckmäßig über einen kleinen Netztransformator mit nachfolgender Gleichrichtung und Glättung bereitgestellt.

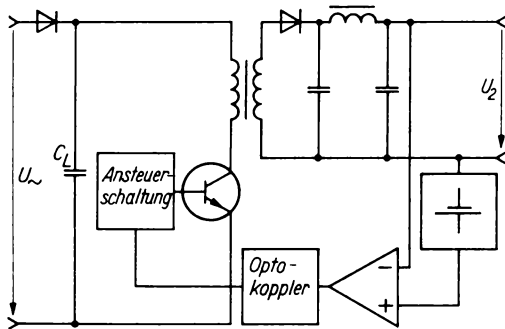


Bild 1.63. Prinzip eines Schaltnetzteils mit Schaltreglerfunktion

Zusammenfassung

Einen besonders hohen Wirkungsgrad bei gleichzeitiger Einsparung des Netztransformators ermöglicht das Schaltnetzteil, das auch gleichzeitig eine Schaltreglerfunktion haben kann. Bei erforderlicher galvanischer Trennung der Gleichspannung vom Netz sind Optokoppler oder zusätzliche Übertrager zur Realisierung der Schaltreglerfunktion erforderlich.

1.7.

Sonderprobleme bei Stromversorgungsschaltungen im Schaltbetrieb

Abschließend sollen einige wichtige Randprobleme von im Schaltbetrieb arbeitenden Stromversorgungsschaltungen betrachtet werden. Beim Einsatz von Elektrolytkondensatoren zur sekundären Siebung entstehen hierbei besondere Schwierigkeiten. Der Ersatzschaltplan des Elektrolytkondensators (Bild 1.64) zeigt, daß neben der bei höheren Frequenzen störenden *Eigeninduktivität* auch ein *Serienwiderstand* zu beachten ist. Dieser entsteht durch den Widerstand des Elektrolyten und steigt mit sinkender Temperatur. Als Folge tritt beim Laden und Entladen ein Längsspannungsabfall auf, der die Ladestromspitzen stark begrenzt, damit den Stromflußwinkel erhöht und eine scheinbare Erniedrigung der wirksamen Kapazität zur Folge hat. Diese Erscheinung tritt bei niedrigen Einsatztemperaturen besonders stark auf. Eine unerwünschte Erwärmung des Kondensators tritt als weitere Folge auf. Da künftig die Schaltfrequenzen weiter steigen werden, sind spezielle Kondensatoren für dieses Einsatzgebiet notwendig. Man kann aber auch statt eines Kondensators mit hoher Kapazität mehrere Kondensatoren mit niedriger Kapazität parallel schalten. Der dann wirksame Serienwiderstand ergibt sich aus der Parallelschaltung der Serienwiderstände der einzelnen Kondensatoren und ist wesentlich niedriger. Zur Verbesserung des Hochfrequenzverhaltens des Siebkondensators sollte zusätzlich ein induktionsarmer Kondensator (z. B. keramischer HDK-Kondensator) parallelgeschaltet werden.

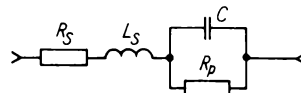


Bild 1.64. Ersatzschaltplan eines Elektrolytkondensators

Der mechanische Aufbau des Kondensators muß impulfest sein, damit durch die impulsmäßig in Abhängigkeit von den Stromspitzen auftretenden Kräfte im Inneren des Wickels keine Beschädigung eintritt. Mehrfachkontaktierungen innerhalb spezieller Kondensatoren erhöhen die Impulsfestigkeit und erniedrigen die Serieninduktivität.

Der *Funkentstörung* von im Schaltbetrieb arbei-

tenden Stromversorgungsteilen ist ebenfalls große Aufmerksamkeit zu widmen. Neben der Schaltfrequenz treten ganzzahlige Vielfache derselben als Störfrequenzen auf, die weit in das Hochfrequenzgebiet reichen (nähere Erläuterungen zu Oberwellen im Abschn. 3.). Rückwirkungen auf das speisende Netz unterdrückt man durch Einsatz von *Netzfiltern* in der Eingangsleitung des Stromversorgungsteils. Diese Netzfilter haben Tiefpaßcharakter und werden meist aus Längsinduktivitäten und Parallelkapazitäten aufgebaut. Eine Abstrahlung von Störspannungen muß durch einen geeigneten Aufbau der Schaltung und des Gehäuses vermieden werden. Besondere Bedeutung hat dabei auch die Wahl eines geeigneten gemeinsamen Massepunktes in der Schaltung. In Übertragern werden zweckmäßig Schirmwicklungen zwischen Primär- und Sekundärwicklung angeordnet.

1.8. Aufgaben

- Unter welchen Betriebsbedingungen wird man die einzelnen Gleichrichterschaltungen sinnvoll einsetzen?
- Analysieren Sie verschiedene Industrieschaltpläne (von Fernsehgeräten, Meßgeräten usw.) nach der Art der verwendeten Gleichrichterschaltung!
- Welche Besonderheiten sind beim Einsatz von Siliziumgleichrichtern zu beachten?
- Konstruieren Sie das Liniendiagramm der Gleichspannung nach Drehstrom-Einweggleichrichtung!
- Begründen Sie, warum eine Anordnung der Sicherung im Lastkreis bei Gegentakt- und Brückenschaltung ungünstig ist!
- Warum sind Verdopplerschaltungen nur bei geringen Lastströmen sinnvoll einsetzbar?
- Erklären Sie die Wirkungsweise der Glättung durch einen Ladekondensator!
- Warum ist der Kapazitätswert des Ladekondensators nach oben begrenzt?
- Welche Bedeutung hat der Stromflußwinkel?
- Wie wirkt ein Siebglied?
- Begründen Sie die Bedingungen für die näherungsweise Berechnung des Siebfaktors!
- Welche Vor- und Nachteile hat ein *RC*-Siebglied?
- Leiten Sie den Gesamtsiebfaktor einer dreigliedrigen Siebkette aus den Einzelsiebfaktoren ab!
- Wie groß ist die Restwechselspannung hinter einem Einweggleichrichter ($U = 30 \text{ V}$) bei einer Stromentnahme von 100 mA und einem Ladekondensator von $1000 \mu\text{F}$?
- Ein Einweggleichrichter mit nachgeschaltetem Ladekondensator und *LC*-Siebglied soll am Ausgang des Siebglieds eine Gleichspannung von 30 V bei einer Stromentnahme von 250 mA liefern. Folgende Werte sind gegeben: $C_L = C_S = 200 \mu\text{F}$; $L = 12 \text{ H}$ mit $R = 100 \Omega$. Berechnen Sie die erforderliche Gleichspannung am Ladekondensator und die Restwechselspannung am Siebkondensator!
- Eine mehrgliedrige Siebkette hinter einer Einweggleichrichtung besteht aus folgenden Bauelementen:
 $L_1 = 15 \text{ H}$; $C_{S1} = 20 \mu\text{F}$; $R_2 = R_3 = 680 \Omega$;
 $C_{S2} = C_{S3} = 50 \mu\text{F}$.
Berechnen Sie den Gesamtsiebfaktor!
- Erklären Sie grafisch und physikalisch die Wirkungsweise der Spannungsstabilisierung mittels *Z*-Diode!
- Weisen Sie in der grafischen Darstellung nach, daß für einen größer werdenden Vorwiderstand bei gleichzeitiger Vergrößerung der Betriebsspannung die stabilisierende Wirkung besser wird!
- Wie wird die resultierende Kennlinie einer Parallelschaltung sowie einer Reihenschaltung zweier Widerstände bestimmt?
- Wie wird aus einer Widerstandsgeraden für den Vorwiderstand der Wert des Vorwiderstands bestimmt?
- Dimensionieren Sie für ein Labornetzgerät eine Stabilisierung für eine Spannung von 6 V bei einem Laststrom zwischen 0 und 250 mA und einer relativen Eingangsspannungsschwankung von $\pm 15 \%$! Prüfen Sie, ob auch ein Widerstand mit $\pm 10 \%$ eingesetzt werden kann!
- Wodurch unterscheiden sich Reihen- und Parallelregelung?
- Weisen Sie nach, daß die Ersatzspannungsquelle bei Spannungskonstanthaltung einen sehr geringen und bei Stromkonstanthaltung einen sehr hohen Innenwiderstand haben muß!
- Erklären Sie an einer Grundschialtung die Wirkungsweise der Spannungs- und Stromkonstanthaltung!
- Unter welchen Betriebsbedingungen tritt die maximale Verlustleistung am Stelltransistor bei der Reihenregelung und bei der Parallelregelung auf?
- Warum ist bei der Reihenregelung eine elektronische Sicherung erforderlich?
- Wie schützt die elektronische Sicherung den Stelltransistor vor Überlastung?
- Welche Vor- und Nachteile haben Festspannungsregler?
- Erklären Sie die Wirkungsweise eines Schaltreglers!
- Wodurch wird die Verlustleistung am Längstransistor eines stetigen Reglers und am Schalttransistor eines Schaltreglers bestimmt?

31. Wie schützt man sich vor ausgangsseitiger Überspannung bei Kurzschluß des Längstransistors?
32. Erklären Sie das Grundprinzip eines Transverters!
33. Worin unterscheiden sich die verschiedenen Transverterarten?
34. Beschreiben Sie die Wirkungsweise des Sperrwandlers!
35. Warum darf der Sperrwandler nicht im Leerlauf betrieben werden?
36. Begründen Sie die im Bild 1.50 gezeigten Strom- und Spannungsverläufe!
37. Wie kann man den Sperrwandler bei Lastausfall vor Überspannungen schützen?
38. Worin besteht der Unterschied zwischen Sperrwandler und Summierwandler?
39. Unter welchen Bedingungen wird der Stromflußwandler eingesetzt?
40. Beschreiben Sie die Wirkungsweise des Gegentaktwandlers!
41. Stellen Sie die Eigenschaften der drei besprochenen Wandlerschaltungen gegenüber!
42. Beschreiben Sie den prinzipiellen Aufbau eines Schaltnetzteils!
43. Welche Vorteile hat ein Schaltnetzteil?
44. Wie erreicht man die galvanische Trennung zwischen Netzwechselspannung und geregelter Gleichspannung im Schaltnetzteil mit Schaltreglerfunktion?

2.

Schaltungen zur Verminderung der Induktionsspannung bei plötzlichen Stromänderungen in Spulen

2.1.

Allgemeines

Im Abschn. 1.4.3. wurde im Bild 1.50 bereits der Strom- und Spannungsverlauf beim Schalten einer Spule an Gleichspannung wiederholend betrachtet. Voraussetzung für diese Darstellung war, daß unmittelbar nach dem Abschalten die Spule über den Widerstand kurzgeschlossen wird. Das ist in der Praxis beim Schalten von Spulen selten der Fall. Meist wird nur der Stromkreis geöffnet. Die im Magnetfeld gespeicherte Energie wird dann an den Spulen- und Schaltkapazitäten in elektrische Feldenergie gewandelt. Nach

$$W = \frac{C}{2} U^2 \quad (2.1)$$

folgt nach Auflösung

$$U = \sqrt{\frac{2W}{C}} \quad (2.2)$$

Für eine kleine Kapazität entsteht folglich eine sehr hohe Selbstinduktionsspannung. Diese führt an Schaltern zum *Funkenüberschlag* zwischen den Kontakten. Die Funkenbildung beansprucht das Kontaktmaterial stark und verursacht eine ungenügende Kontaktgabe (Fein- und Grobwanderung an den Kontakten) bis zum völligen Ausfall der Kontakte. Unter ungünstigen Bedingungen kann beim Schalten einer Spule an Gleichspannung eine Lichtbogenbildung zwischen den Kontakten eintreten (Zerstörung der Kontakte).

Auch ohne unmittelbare Funkenbildung kann die Funktion elektronischer Geräte gestört werden. So kann an als Schalter verwendeten Halbleiterbauelementen die hohe Selbstinduktionsspannung zur Zerstörung führen (s. auch Abschn. 5.2.).

2.2.

Schaltungsvarianten

Ursache für die Induktionsspannung ist die Umwandlung der magnetischen Feldenergie in elektrische Feldenergie an den Wicklungs- und Schaltkapazitäten beim Unterbrechen des Stromflusses. Erhöht man die wirksame Kapazität durch *Parallelschalten eines Kondensators* zur Schaltstrecke nach Bild 2.1, so wird die Größe der Selbstinduktionsspannung vermindert. Die Gesamtenergie verbraucht sich in Form einer gedämpften Schwingung oder einer aperiodischen Entladung. Die Kapazität darf dabei nicht zu groß sein, da sonst bei einem rasch erfolgenden erneuten Einschalten die vom Kondensator aufgenommene Energie noch nicht wieder abgebaut wäre und vom Schalter kurzgeschlossen würde. Die Wirkung dieses Kurzschlusses kann man durch die *Reihenschaltung eines Widerstands* mit dem Kondensator nach Bild 2.2 mindern.

Wird nach Bild 2.3 lediglich ein *Widerstand* zur Schaltstrecke parallelgeschaltet, so erreicht man ebenfalls eine Minderung der auftretenden

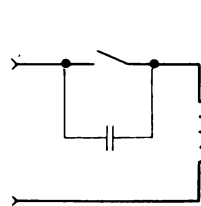


Bild 2.1. Funkenminderung durch Parallelkondensator

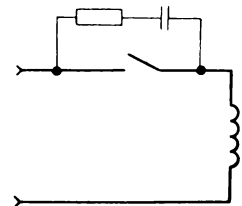


Bild 2.2. Funkenminderung durch Kondensator mit Widerstand

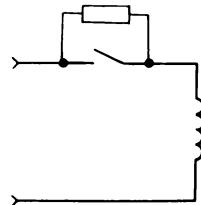


Bild 2.3. Funkenminderung durch Parallelwiderstand

Spannungsspitze. Da dieser Widerstand aus Funktionsgründen nicht zu niedrig sein darf, ist seine induktionsspannungsbegrenzende Wirkung gering.

Eine besonders gute Verminderung der Induktionsspannung wird durch Parallelschalten einer Diode zur zu schaltenden Spule erreicht (Bild 2.4). Die Diode muß dabei so gepolt sein, daß sie für die Betriebsspannung in Sperrrichtung liegt. Im Bild 2.5 a sind der Stromfluß und die Polarität des Spannungsabfalls für den Einschaltfall gekennzeichnet. Wird jetzt der Schaltkontakt geöffnet, so wird nach dem Lenzschen Gesetz der durch die Selbstinduktionsspannung angetriebene Strom die ursprüngliche Richtung beibehalten. Die Spule wirkt jetzt wie eine Spannungsquelle mit der im Bild 2.5 b eingetragenen Polarität. (Innerhalb der Spannungsquelle fließt der Strom von $-$ nach $+$.) Die Diode stellt für diese Selbstinduktionsspannung einen Kurzschluß dar. Der die Diode durchfließende Spitzenwert des Stromes ist dabei gleich dem Strom, der im Abschaltmoment die Spule durchfließt.

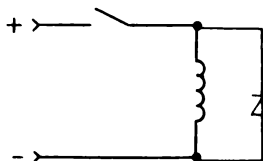
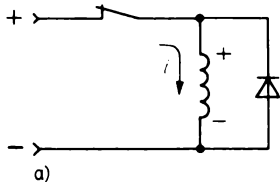
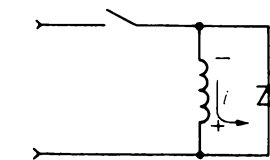


Bild 2.4. Funkenminde-
rung durch Diode



a)



b)

Bild 2.5. Richtung
von Strom und Span-
nung
a) während der Ein-
schaltzeit;
b) im Abschaltmo-
ment

Die Schaltungen mit Kondensator und Diode sind nur beim Schalten an Gleichspannung einsetzbar. Es tritt jedoch auch beim Schalten an Wechselfspannung je nach Phasenlage von Strom und Spannung im Zeitpunkt des Schaltens eine Induktionsspannung auf. Durch Parallelschalten eines Varistors zur Spule nach Bild 2.6 kann man auch hier Abhilfe schaffen. Voraussetzung ist hierbei, daß die Sperrspannung des Varistors größer als der Spitzenwert der angelegten Wechselfspannung ist.

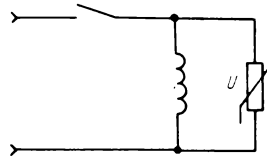


Bild 2.6. Funkenmin-
derung durch Varistor

Zusammenfassung

Beim Unterbrechen des Stromflusses in Stromkreisen mit Spulen entsteht eine hohe Induktionsspannung, die den Schalter stark beansprucht. Die Induktionsspannung wird vermindert, indem man die im magnetischen Feld gespeicherte Energie mehr oder weniger kurzschließt oder sie in einem Kondensator in elektrische Feldenergie wandelt.

2.3.

Aufgaben

1. Begründen Sie, warum es bei einer Spule zu einer hohen Induktionsspannung beim Abschalten kommt!
2. Welche Folgen hat die Funkenbildung an Kontakten?
(Welche Werkstoffe sind hiergegen besonders empfindlich?)
3. Wie entsteht die induktionsspannungsbegrenzende Wirkung der genannten Schaltungen?
4. Warum ist die Schaltung mit Kondensator und mit Diode nur beim Schalten von Gleichspannung einsetzbar?

3.

Verstärker mit elektronischen Bauelementen

3.1.

Verstärker als Vierpol

Ein Verstärker ist eine Schaltung von elektronischen Bauelementen mit zwei Eingangsbuchsen (1 und 1') und zwei Ausgangsbuchsen (2 und 2'). Durch ihn wird durch Hinzufügen von Energie aus einer äußeren Quelle eine am Eingang angelegte elektrische Wechselgröße (Spannung, Strom, Leistung) oder sich langsam ändernde Gleichgröße vergrößert (Bild 3.1).

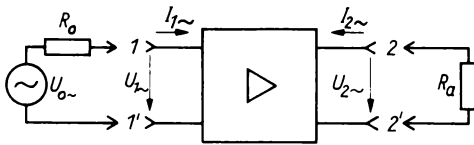


Bild 3.1. Verstärker, allgemein

Die Verstärker können nach verschiedenen Gesichtspunkten eingeteilt werden, je nachdem, ob man ihren Aufbau, ihre Wirkungsweise oder ihre Verwendung besonders charakterisieren will (Tafel 3.1).

Im Vorverstärker muß eine im allgemeinen sehr kleine Steuerleistung $P_{1\sim} = U_{1\sim} I_{1\sim}$ auf eine Wechsellleistung $P_{2\sim} = U_{2\sim} I_{2\sim}$ verstärkt werden, die wesentlich größer gegenüber möglichen Störleistungen ist und eine unproblematische Signalverarbeitung bzw. Weiterverstärkung gestattet.

Da die Wechselgrößen noch relativ klein sind, können die für den Verstärkervorgang maßgebenden Transistorkennlinien im Aussteuerungsbereich als geradlinig betrachtet werden, und man kann mit einer verzerrungsarmen Übertragung rechnen. Wegen der relativ geringen Ausgangsleistung interessiert bei diesen Verstärkern der Wirkungsgrad nicht. Er kann sehr klein sein.

Im Endverstärker wird die vom Vorverstärker abgegebene Leistung auf den erforderlichen Endwert verstärkt, wobei im allgemeinen noch keine Verzerrungen durch Übersteuern auftreten sollen. Dazu ist die maximal abzugebende

verzerrungsfreie Leistung eines Transistors zu bestimmen. Der Verstärkungsgrad spielt eine untergeordnete Rolle, da die Steuerleistung durch den Vorverstärker bereitgestellt wird. Die Ausgangsleistung liegt zwischen der des Vor- und der des Senderverstärkers. Der Wirkungsgrad ist bei stationären, netzbetriebenen Anlagen zweitrangig.

Beim Senderverstärker kann die Ausgangsleistung mehrere Kilowatt betragen; damit muß der Wirkungsgrad, d. h. das Verhältnis von abgegebener Nutz- zu aufgewandter Gesamtleistung, sehr groß sein. Der Verstärkungsgrad und die Verzerrungsfreiheit sind diesem Problem unterzuordnen. Ein hoher Wirkungsgrad ist nicht nur wegen geringer Energiekosten anzustreben, sondern auch, weil die nicht in Wechsellleistung umgesetzte Gleichleistung als Erwärmung des Verstärkerelements auftritt und seine Belastbarkeit begrenzt.

Tafel 3.1. Einteilung der Verstärker

Einteilung nach	Verstärkerarten
Bauelement	Röhrenverstärker, Transistorverstärker
Kopplungsart	direktgekoppelter Verstärker, RC-gekoppelter Verstärker, transformatorgekoppelter Verstärker, bandfiltergekoppelter Verstärker, diodegekoppelter Verstärker, optoelektronisch-gekoppelter Verstärker
Arbeitsweise	Eintaktverstärker, Gegentaktverstärker
Betriebsart	A-Verstärker, B-Verstärker, AB-Verstärker, C-Verstärker
Frequenzband	Selektivverstärker, Breitbandverstärker
Arbeitsfrequenz	Gleichspannungsverstärker, Impulsverstärker, Niederfrequenzverstärker, Hochfrequenzverstärker
Leistung	Vorverstärker, Endverstärker, Senderverstärker

Kenndaten eines Verstärkers

Verstärkungsgrad oder Verstärkung v

Leistungsverstärkung

$$v_p = \frac{P_{2\sim}}{P_{1\sim}} = \frac{U_{2\sim} I_{2\sim}}{U_{1\sim} I_{1\sim}} = v_u v_i \quad (3.1)$$

mit Spannungsverstärkung $v_u = \frac{U_{2\sim}}{U_{1\sim}} \quad (3.2)$

und Stromverstärkung $v_i = \frac{I_{2\sim}}{I_{1\sim}} \quad (3.3)$

Das Verhältnis der Eingangs- bzw. Ausgangsgrößen liefert

Eingangswiderstand $R_{\text{ein}} = \frac{U_{1\sim}}{I_{1\sim}} \quad (3.4)$

Ausgangswiderstand $R_{\text{aus}} = \frac{U_{21\sim}}{I_{2k\sim}} \quad (3.5)$

$U_{21\sim}$ Leerlaufspannung ($R_a \rightarrow \infty$)

$I_{2k\sim}$ Kurzschlußstrom ($R_a = 0$)

Arbeitswiderstand $R_a = \frac{U_{2\sim}}{I_{2\sim}} \quad (3.6)$

(Lastwiderstand)

$$\frac{R_a}{R_{\text{ein}}} = \frac{U_{2\sim}/I_{2\sim}}{U_{1\sim}/I_{1\sim}} = \frac{U_{2\sim}}{U_{1\sim}} \cdot \frac{1}{I_{2\sim}/I_{1\sim}} = \frac{v_u}{v_i}$$

Die Wechselgrößen können dabei einen beliebigen, also auch nichtsinusförmigen Kurvenverlauf haben.

Oft gibt man Spannungs-, Strom- und Leistungsverhältnisse im logarithmischen Maß an, weil es Rechenvorteile bringt.

$$v_{p/\text{dB}} = 10 \lg \frac{P_{2\sim}}{P_{1\sim}}$$

$$v_{u/\text{dB}} = 20 \lg \frac{U_{2\sim}}{U_{1\sim}}$$

$$v_{i/\text{dB}} = 20 \lg \frac{I_{2\sim}}{I_{1\sim}}$$

Frequenzbandbreite b

Frequenzbandbreite $b = f_o - f_u \quad (3.7)$

f_o obere Grenzfrequenz

f_u untere Grenzfrequenz.

Die Grenzfrequenz ist im allgemeinen die Frequenz, bei der die Verstärkung um 3 dB auf $1/\sqrt{2}$ gegenüber der maximalen Verstärkung abgesunken ist, also

$$\frac{v_u}{v_{u\text{max}}} = \frac{1}{\sqrt{2}} = 0,707$$

Nichtlineare Verzerrungen

Nichtlineare Verzerrungen werden bei Verstärkern am häufigsten durch den Klirrfaktor k angegeben.

Klirrfaktor $k = \frac{\text{Effektivwert aller Oberwellen}}{\text{Effektivwert aller Harmonischen}}$

$$k = \frac{\sqrt{U_{2\sim\text{eff}}^2 + U_{3\sim\text{eff}}^2 + \dots}}{\sqrt{U_{1\sim\text{eff}}^2 + U_{2\sim\text{eff}}^2 + U_{3\sim\text{eff}}^2 + \dots}} \quad (3.8)$$

Der Klirrfaktor wird meist in % angegeben. Die Entstehung des Klirrfaktors wird im Abschn. 3.6.1. näher erläutert.

Störabstand

Der Störabstand b_{st} ist das Verhältnis von Nutz- zu Störgröße an derselben Stelle, meist am Verstärkerausgang, und kann als Spannungs- oder Leistungsverhältnis angegeben werden.

$$b_{\text{st}} = \left| \frac{U_{2\sim\text{Nutz}}}{U_{2\sim\text{Stör}}} \right|$$

Überschwingen

Durch starke Impulse können die im Verstärker durch die Bauelemente und deren Aufbau ungewollt entstandenen Schwingkreise zur Ausbildung einer gedämpften Schwingung angestoßen werden. Man untersucht solche Überschwingvorgänge, indem man am Eingang eine Rechteckspannung anlegt und die Ausgangsspannung oszilloskopiert.

Brummodulation

Die von ungenügender Siebung der Netzbrummspannung sowie von Einstreuungen (fehlerhafte Abschirmung, ungünstige Anordnung der Bauelemente) im Ausgangskreis der Endverstärker herrührende Brummodulation m_{Br} ist eine Störmodulation. Sie ist als Verhältnis der Amplitudenänderung zur mittleren Amplitude bei konstanter Eingangsspannung definiert.

$$m_{\text{Br}} = \frac{\Delta \hat{U}_2}{\hat{U}_{2\text{mittel}}} \quad U_1 = \text{konst.} \quad (3.9)$$

3.2.

Grundschaltungen der Verstärker

3.2.1.

Allgemeines

Das Verstärkerbauelement Transistor kann auf verschiedene Arten in den allgemeinen Vierpol (s. Bild 3.1) eingesetzt werden.

Die Verstärkergrundschaltungen werden nach der dem Ein- und Ausgang gemeinsamen Elektrode des Verstärkerbauelements benannt. Alle Schaltungsarten haben voneinander abweichende Eigenschaften, die man zum Lösen der vielfältigen Aufgaben kennen muß.

Schaltungen mit bipolaren Transistoren können zu den entsprechenden Schaltungen mit unipolaren Transistoren (Feldeffekttransistoren) analog betrachtet werden. Während bei letzteren im allgemeinen nur drei Größen (Steuerspannung, Ausgangsspannung und Ausgangsstrom) eine Rolle spielen, sind bei Verstärkern mit bipolaren Transistoren wegen der notwendigen Leistungssteuerung jedoch vier Größen, nämlich die Ströme und Spannungen an Ein- und Ausgang, in Rechnung zu setzen. Die folgenden Ausführungen beschränken sich, bis auf die Ausführungen zur „Kombination Unipolar- und Bipolartransistor“, auf Schaltungen mit bipolaren Transistoren.

Die im Abschn. 3.2.3. behandelten Schaltungen mit mehreren Transistoren sind erweiterte Grundschaltungen, um speziellere Forderungen erfüllen zu können.

3.2.2.

Grundschaltungen mit einem Transistor

Emitterschaltung

Diese Grundschaltung findet in allen Bereichen der Elektronik die vielfältigsten Anwendungsmöglichkeiten, da mit ihr eine große Leistungsverstärkung erzielt werden kann. Nachteilig ist, daß im Vergleich zur Basisschaltung nur eine kleinere obere Grenzfrequenz erreichbar ist.

Gemeinsame Bezugselektrode für Eingangs- und Ausgangskreis ist der Emitter. Durch einen kleinen Eingangsstrom I_{B-} wird der Ausgangsstrom I_{C-} gesteuert. Bei Erhöhung der Eingangsspannung wird der Eingangs- und damit der Ausgangsstrom vergrößert. Da dadurch der Spannungsabfall über dem Arbeitswiderstand steigt, sinkt bei konstanter Betriebsspannung die Ausgangsspannung zwischen Kollektor und

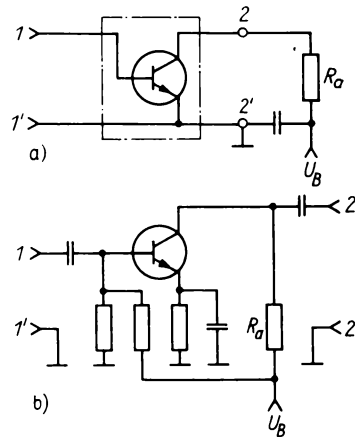


Bild 3.2. Emittterstufe

a) Prinzip; b) Schaltungsbeispiel

Emitter, ist also gegenphasig zur Eingangsspannung (Bild 3.2).

Neben der Spannungsverstärkung findet auch eine Strom- und Leistungsverstärkung statt.

Die Eingangs- und die Ausgangswiderstände liegen im mittleren Wertebereich, unterscheiden sich untereinander wesentlich weniger als bei den anderen Grundschaltungen und sind von den Abschlußverhältnissen nur gering abhängig. Als Richtwerte gelten

$$R_{\text{aus}} \approx 20 \dots 100 \text{ k}\Omega$$

$$R_{\text{ein}} \approx 500 \Omega \dots 2 \text{ k}\Omega.$$

Bei mehrstufigen Verstärkern mit großem Verstärkungsgrad wendet man meistens diese Schaltung an.

Basisschaltung

Die Basisschaltung wird angewendet, wenn ein kleiner Eingangs- und ein großer Ausgangswiderstand sowie eine hohe obere Grenzfrequenz verlangt wird, z. B. bei UKW-Eingangsstufen.

Gemeinsame Bezugselektrode für Eingangs- und Ausgangskreis ist die Basis. Da diese zwischen Emitter und Kollektor liegt, wird die bei höheren Frequenzen schädliche Rückwirkungsgröße zwischen Eingang und Ausgang vermindert.

Die Steuerung erfolgt durch den Eingangsstrom, der stets etwas größer als der Ausgangsstrom ist. Die Stromverstärkung liegt damit geringfügig unter dem Wert 1; dagegen erzielt man eine hohe Spannungsverstärkung (Bild 3.3).

Eine positive Eingangsspannung hebt das Emitterpotential an, verringert den Ausgangsstrom

und damit auch den Spannungsabfall über dem Arbeitswiderstand. Die Ausgangswechselspannung zwischen Kollektor und Basis steigt dadurch bei konstanter Betriebsspannung an, ist also gleichphasig zur Eingangsspannung.

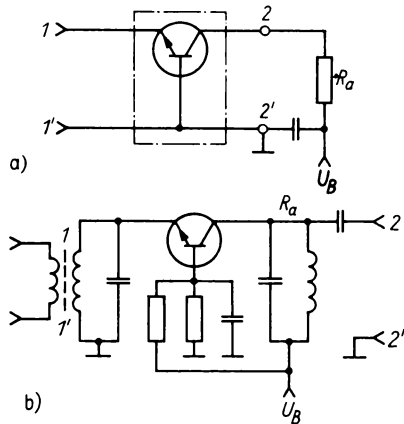


Bild 3.3. Basisstufe

a) Prinzip; b) Schaltungsbeispiel

Wegen des im Eingangskreis fließenden hohen Emittersstroms und der ohnehin niederohmigen Basis-Emitter-Strecke ergibt die Basisschaltung einen sehr kleinen Eingangswiderstand von etwa

$$R_{\text{ein}} \approx 30 \, \Omega \dots 1 \, \text{k}\Omega.$$

Dagegen liegt der Ausgangswiderstand bei

$$R_{\text{aus}} \approx 100 \, \text{k}\Omega \dots 1 \, \text{M}\Omega.$$

Kollektorschaltung

Die Kollektorschaltung hat einen hohen Eingangs- und einen niedrigen Ausgangswiderstand. Sie wird deshalb als Impedanzwandler eingesetzt, wenn hochohmige Quellen an niederohmige Verbraucher angepaßt werden müssen.

Gemeinsame Bezugselektrode für Eingangs- und Ausgangskreis ist der Kollektor. Der Arbeitswiderstand liegt zwischen Emitter und Bezugspotential. Die Eingangsspannung wird zwischen Basis und Kollektor gelegt und durch die an R_a abfallende Ausgangsspannung 100%ig gegengekoppelt.

Liegt die positive Halbwelle am Eingang, dann steigt der Ausgangsstrom und damit die Ausgangsspannung; sie ist also zur Eingangsspannung gleichphasig. Gleichzeitig folgt aber die Emitterspannung bis auf eine geringe Differenz der Eingangsspannung. (Die Schaltung heißt

deshalb auch Emitterfolger.) Die Spannungsverstärkung liegt damit unter dem Wert 1; eine Strom- und Leistungsverstärkung findet aber statt.

Da zwischen Basis und Emitter nur die aus Eingangs- und Ausgangsspannung gebildete Differenzspannung wirksam ist, müßte eine wesentlich höhere Eingangsspannung als bei der Emitterschaltung aufgebracht werden, um die gleiche Steuerwirkung zu erreichen. Das ist gleichbedeutend einem großen Eingangswiderstand. Die Ausgangsspannung wird nahezu unabhängig von der Größe des Arbeitswiderstands nur von der geringfügig größeren Eingangsspannung bestimmt. Quellen, deren Klemmenspannung aber belastungsunabhängig sind, haben einen kleinen Innenwiderstand. Da der Ausgangswiderstand vom Innenwiderstand bestimmt wird, muß also R_{aus} der Kollektorschaltung klein sein (Bild 3.4). Als Richtwerte gelten

$$R_{\text{ein}} \approx 3 \, \text{k}\Omega \dots 1 \, \text{M}\Omega$$

$$R_{\text{aus}} \approx 30 \, \Omega \dots 1 \, \text{k}\Omega.$$

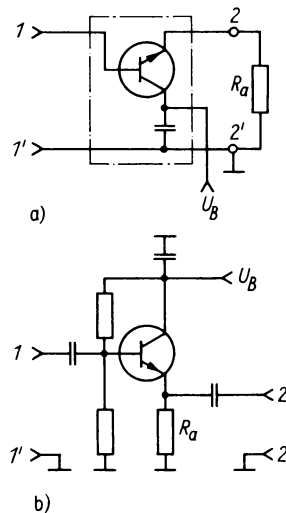


Bild 3.4. Kollektorstufe

a) Prinzip; b) Schaltungsbeispiel

Bootstrap-Schaltung

Die Bootstrap-Schaltung arbeitet nach einem Schaltungsprinzip, bei dem die Spannungen zweier Punkte einer Schaltung „mitlaufen“. Dadurch bleibt zwischen diesen beiden Punkten die Spannungsdifferenz zeitlich konstant, und es kann sich keine Wechselspannung ausbilden. Das Bootstrap-Prinzip wird u. a. angewendet, um einen Widerstand dynamisch zu vergrößern (s. Abschn. 3.6.3.) oder um die von einem Generator erzeugte Sägezahnspannung zu linearisieren (s. Abschn. 4.3.3.).

Im vorliegenden Fall wird das Bootstrap-Prinzip in einer Kollektorschaltung angewendet, um den Eingangswiderstand zu erhöhen. Wird nämlich entsprechend Bild 3.4 die Basisvorspannung über einen Basisspannteiler gewonnen, dann kann der prinzipiell mögliche Eingangswiderstand nicht erreicht werden. Auf den Eingang bezogen, liegen die Teilerwiderstände parallel. Der Widerstand dieser Parallelschaltung ist meist kleiner als der Transistoreingangswiderstand und damit bestimmend für den Gesamteingangswiderstand. Der Einfluß des Spannungsteilers läßt sich teilweise eliminieren, indem die an ihm abgegriffene Vorspannung über den Widerstand R der Basis zugeführt wird. Außerdem wird R kapazitiv über den Kondensator C mit dem Ausgang verbunden (Bild 3.5). Da die Ausgangsspannung U_2 nur geringfügig kleiner als die Eingangsspannung U_1 und mit dieser gleichphasig ist, ist der

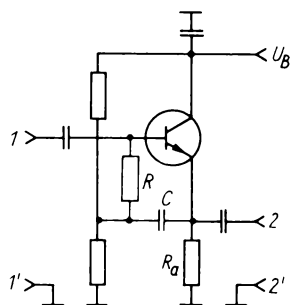


Bild 3.5. Kollektorstufe mit Bootstrap-schaltung

Wechselspannungsabfall über R als Differenz von U_1 und U_2 sehr klein. Ein sehr kleiner Spannungsabfall läßt auf einen sehr kleinen Strom schließen, d. h., der Widerstand R erscheint dynamisch vergrößert und damit auch der Eingangswiderstand der Schaltung. Zusammenfassend werden in Tafel 3.2 die Eigenschaften der Emitter-, Basis- und Kollektorschaltung wiedergegeben. Die angeführten Eingangs- und Ausgangswiderstände liegen meist innerhalb der angegebenen Bereiche. Sie sind stark abhängig vom Bauelementtyp und seinen Parametern und der angeschlossenen Schaltung. So ist der Eingangswiderstand stark abhängig vom Arbeitswiderstand, der zwischen Null (Kurzschluß) und unendlich (Leerlauf) liegen kann. Der Ausgangswiderstand ist dagegen vom Innenwiderstand R_0 der Steuerspannungsquelle abhängig.

3.2.3.

Grundsaltungen mit mehreren Transistoren

Kaskadenschaltung (Darlington-Schaltung)

Bei einer von Darlington angegebenen Kaskadenschaltung werden zwei, höchstens drei Transistoren so zusammengeschaltet, daß der Emitter- bzw. Kollektorstrom des vorhergehenden Transistors ist.

Tafel 3.2. Zusammenstellung der wichtigsten Eigenschaften der Emitter-, Basis- und Kollektorschaltung

Emitterschaltung	Basisschaltung	Kollektorschaltung
Eingangswiderstand R_{ein} 500 Ω ...2 k Ω	30 Ω ...1 k Ω	3 k Ω ...1 M Ω
Ausgangswiderstand R_{aus} 20...100 k Ω	100 k Ω ...1 M Ω	30 Ω ...1 k Ω
Spannungsverstärkung v_u max. $10^4 \gg 1$	max. $10^4 \gg 1$	< 1
Stromverstärkung v_i max. $10^2 \gg 1$	< 1	max. $10^2 \gg 1$
Leistungsverstärkung v_p max. 10^6	max. 10^4	max. 10^2
Phasenwinkel der Spannungsverstärkung bei reellem Arbeitswiderstand 180°	0°	0°
Bevorzugte Anwendung bei mehrstufigen Verstärkern mit großem Verstärkungsgrad	bei rückwirkungsfreien Verstärkern in UKW-Eingangsstufen; bei Verstärkern mit einer sehr hohen oberen Grenzfrequenz; bei Endstufen in UKW-Sendern	als Impedanzwandler zur Anpassung hochohmiger Quellen an niederohmige Verbraucher

Mit dieser Anordnung erhält man eine hohe Stromverstärkung, die sich etwa aus dem Produkt der Stromverstärkungen der Einzeltransistoren ergibt.

Da sich dabei auch der Eingangswiderstand wesentlich erhöht, ist die Kaskadenschaltung nicht nur als hochverstärkende Einheit, sondern auch als Trennstufe (Impedanzwandler) und hochohmige Eingangsstufe einsetzbar.

Kaskadenschaltungen lassen sich realisieren als

- Standard-Darlington-Schaltungen mit Transistoren gleichen Leitungstyps (Bild 3.6 a) und als
- Komplementär-Darlington-Schaltungen mit Transistoren unterschiedlichen Leitungstyps (Bild 3.6 b).

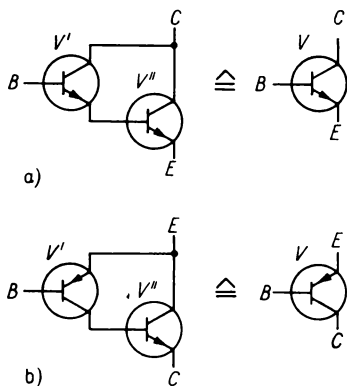


Bild 3.6. Darlingtonschaltungen

a) Standard-Darlingtonschaltung; b) Komplementär-Darlingtonschaltung

Jede Darlington-Schaltung kann durch einen Ersatztransistor dargestellt werden. Bei der Standard-Darlington-Schaltung ist die Zonenfolge des Ersatztransistors die der Einzeltransistoren, bei der Komplementär-Darlington-Schaltung die des ersten Einzeltransistors.

Darlington-Anordnungen sind einfach integrierbar. Mit ihnen können alle Transistorschaltungen aufgebaut werden; ihre Eigenschaften werden durch die des Ersatztransistors bestimmt. Als Beispiel ist im Bild 3.7 eine Emitterstufe wiedergegeben, sowohl mit einer Standard- als auch mit einer Komplementär-Darlington-Schaltung.

Kaskodeschaltung

Die Kaskodeschaltung ist eine 2stufige Eingangsschaltung. Sie wird wegen ihrer geringen dynamischen Eingangskapazität vorwiegend bei hohen Frequenzen eingesetzt.

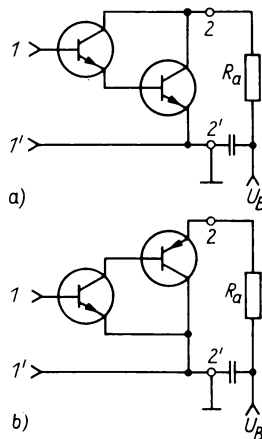


Bild 3.7. Emitterstufe
a) mit Standard-Darlingtonschaltung
b) mit Komplementär-Darlingtonschaltung

Der Transistor $V1$ arbeitet in Emitter-, der Transistor $V2$ in Basisschaltung.

Im Vergleich zu einer einfachen Emitterschaltung weist die Kaskodeschaltung bei tiefen und mittleren Frequenzen etwa gleiches Verhalten aus. Das betrifft sowohl den Eingangswiderstand, der gleich dem der ersten Stufe ist, als auch die Gesamtspannungsverstärkung, die von der zweiten Stufe bestimmt wird. Der Vorteil der wesentlich geringeren dynamischen Eingangskapazität wird erst bei hohen Frequenzen wirksam.

Die unerwünschten Transistor- und Schaltkapazitäten C_1 und C_2 (Bild 3.8) bestimmen die wirksame Eingangskapazität C_{ein} nach der Beziehung

$$C_{\text{ein}} = C_1 + C_2(1 - v_u).$$

Bei der einfachen Emitterschaltung ist $|v_u| \gg 1$, so daß die Eingangskapazität vor allem von $C_2 v_u$ bestimmt wird. Bei der Kaskodeschaltung dagegen wird die erste Stufe durch den kleinen Eingangswiderstand der in Basisschaltung betriebenen zweiten Stufe belastet. Dadurch ist die für C_{ein} bestimmende Spannungsverstärkung der ersten Stufe nur $v_u \approx -1$ (Minus wegen Phasendrehung), und es folgt

$$C_{\text{ein}} = C_1 + 2 C_2.$$

Die Eingangskapazität ist bei der Kaskodeschaltung also wesentlich geringer als bei einer einfachen Emitterstufe mit großem Verstärkungsgrad.

Kaskodeschaltungen können sowohl mit Transistoren des gleichen als auch mit Transistoren ungleichen Leitungstyps aufgebaut werden; sie lassen sich integrieren. Bild 3.8 zeigt zwei einfache Schaltungsbeispiele.

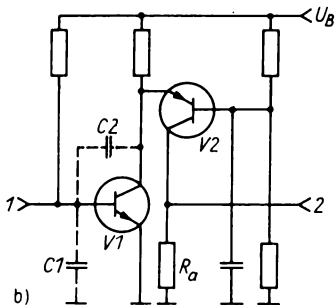
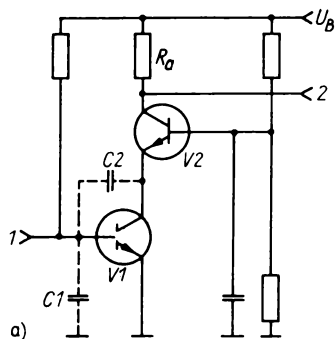


Bild 3.8. Kaskodeschaltung

a) mit Transistoren gleichen Leistungstyps

b) mit Transistoren unterschiedlichen Leistungstyps

Kombination Unipolar- und Bipolartransistor

Soll der Eingangswiderstand einer Schaltung sehr groß sein, empfiehlt es sich, einen Unipolartransistor (Feldeffekttransistor FET) einzusetzen.

Unipolartransistoren lassen sich nahezu leistungslos steuern; den Nachteil ihrer relativ geringen Steilheit (Verhältnis von Ausgangsstrom zu Steuerspannung) kann man durch einen nachgeschalteten Bipolartransistor beseitigen.

Bild 3.9 zeigt die prinzipielle Schaltung einer Eingangsstufe mit sehr hohem Eingangswiderstand. Der Transistor V1 ist ein MOS-Feldefeffekttransistor mit p-Kanal vom Anreicherungstyp. Er arbeitet in Source-, der nachfolgende Bipolartransistor V2 in Kollektorschaltung. Der Eingangswiderstand und die Gesamtverstärkung

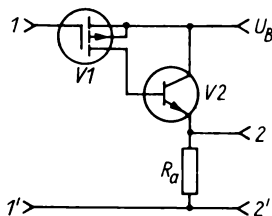


Bild 3.9. Kombination Unipolar- und Bipolartransistor

werden von der ersten Stufe bestimmt, wobei durch V2 die Steilheit des Transistors V1 mit dem Faktor der Kurzschlußstromverstärkung vergrößert wird.

Die Transistorkombination weist außerdem ein günstiges Rauschverhalten aus. Rauschen ist eine unerwünschte Störleistung und wird von der Schaltung verursacht. Bei jeder mehrstufigen Anordnung wird das Gesamttrauschen vor allem von der ersten Stufe bestimmt. Unipolartransistoren haben aber im Vergleich zu Bipolartransistoren bessere Rauschwerte (kleinere Rauschzahl), so daß die beschriebene integrierbare Transistorkombination auch in dieser Beziehung vorteilhaft ist.

Nachteilig ist, daß auf Grund der großen resultierenden Steilheit der Arbeitspunkt sehr sorgfältig stabilisiert werden muß (im Bild 3.9 nicht dargestellt). Dadurch wird eine 5- bis 10mal so hohe Betriebsspannung als sonst für Halbleiterschaltungen üblich benötigt.

Differenzverstärkerschaltung

Die Differenzverstärkerschaltung (DV) besteht aus zwei parallelliegenden, datengleichen Transistorverstärkern in Emitterschaltung (Bild 3.10). Der Strom für die beiden Transistoren wird über eine Konstantstromquelle eingespeist. Die erforderlichen Basis-Gleichströme werden von den anzuschließenden Quellen geliefert.

Eine Konstantstromquelle liefert wegen ihres sehr großen Innenwiderstands ($> 10 \text{ M}\Omega$) einen vom Außenwiderstand fast unabhängigen, konstanten Strom I_0 .

Ohne Aussteuerung fließt durch den Transistor V1 der gleiche Strom wie durch den Transistor

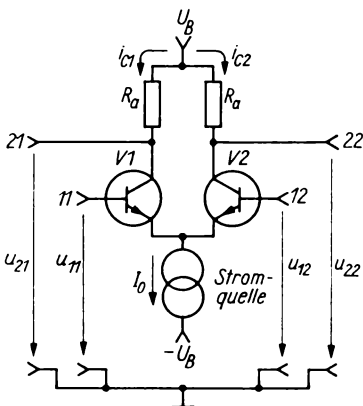


Bild 3.10. Differenzverstärkerschaltung

V_2 , so daß auch die beiden Ausgangsspannungen gleich sind.

Im folgenden werden verschiedene Betriebsfälle für DV beschrieben.

● *Beide Eingangsspannungen ändern sich genau gegenphasig.*

Nimmt man an, daß sich z. B. die Eingangsspannung u_{11} um $+5\text{ mV}$ und die Eingangsspannung u_{12} um -5 mV ändert, dann erfolgt wegen des konstanten Gesamtmitterstroms nur eine Umsteuerung zwischen i_{C1} und i_{C2} . Im gewählten Beispiel wird i_{C1} um den gleichen Betrag größer, wie i_{C2} kleiner wird. Damit sinkt die Ausgangsspannung u_{21} im gleichen Maße, wie die Ausgangsspannung u_{22} steigt.

● *Beide Eingangsspannungen ändern sich in gleicher Weise.*

Dieser Betriebsfall heißt Gleichtaktbetrieb. Er läßt sich leicht verstehen, wenn man sich vorstellt, daß beide Eingänge miteinander verbunden wären. Beide Transistoren sind dann parallel geschaltet und werden durch eine Eingangsspannung in gleicher Weise gesteuert. Die Stromaufteilung durch die Transistoren bleibt dadurch unverändert, so daß sich auch die Ausgangsspannungen nicht ändern. Gleichtaktsignale werden also bei idealen Verhältnissen nicht übertragen.

● *Beide Eingangsspannungen ändern sich unabhängig voneinander.*

Dieser allgemeine Betriebsfall ergibt sich als Überlagerung der beiden oben beschriebenen Fälle. Die Aufteilung des konstanten Gesamtstroms auf die Transistoren V_1 und V_2 hängt von der Differenz der beiden Eingangsspannungen ab. Je größer diese Differenz ist, desto unterschiedlicher sind die beiden Kollektorströme und damit die verstärkten Ausgangsspannungen. Ist die Differenz Null (s. Gleichtaktbetrieb), dann sind die Kollektorströme und beide Ausgangsspannungen gleich groß. Die zu verstärkende Spannung wird der Schaltung als Differenzspannung zugeführt. Das kann auch dadurch geschehen, daß z. B. der Eingang 12 mit Masse verbunden wird und am Eingang 11 nun die Spannung u_{11} als Differenzspannung zwischen 11 und 12 wirkt. Da aber $u_{12} = 0$ (12 ist mit Masse verbunden), ist $u_D = u_{11} - u_{12} = u_{11}$. Im Bild 3.11 sind die Strom- und Spannungsverläufe für willkürlich gewählte Eingangsspannungsverläufe prinzipiell dargestellt.

Differenzverstärkerschaltungen werden, von wenigen Sonderfällen abgesehen, in integrierter

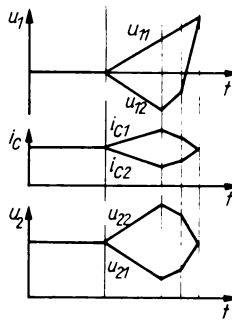


Bild 3.11. Strom- und Spannungsverläufe in der Differenzverstärkerschaltung

Technik ausgeführt. Das hat den Vorteil, daß auf dem gleichen Substrat alle Bauelemente untergebracht sind und die paarigen Elemente in ihren elektrischen Parametern weitgehend übereinstimmen. Außerdem treten bei Betrieb nur unbedeutende Temperaturdifferenzen zwischen den Elementen auf, so daß durch Temperatureinflüsse die guten Symmetrieeigenschaften kaum beeinträchtigt werden.

Kenndaten der Differenzverstärkerschaltung

Differenzverstärkung v_D

$$v_D = \frac{u_{22}}{u_D} = - \frac{u_{21}}{u_D} \quad (3.10)$$

u_D Differenzspannung, $u_D = u_{11} - u_{12}$.

Praktisch erreichbare Werte: $|v_D| \approx 10^3 \hat{=} 60\text{ dB}$.

Gleichtaktverstärkung v_G

$$v_G = \frac{u_{22}}{u_G} = \frac{u_{21}}{u_G} \quad (3.11)$$

u_G Gleichtaktspannung, $u_G = u_{11} = u_{12}$.

Ideal wäre $v_G = 0$; praktisch erreichbare Werte: $v_G \approx 10^{-2} \hat{=} -40\text{ dB}$.

Gleichtaktunterdrückung CMR (common mode rejection)

$$CMR = \left| \frac{v_D}{v_G} \right| \quad (3.12)$$

Ideal wäre $CMR \rightarrow \infty$; praktisch erreichbare Werte: $CMR \approx 10^3 \dots 10^5 \hat{=} 60 \dots 100\text{ dB}$.

Konstantstromquellen und Stromspiegel

Konstantstromquellen liefern wegen ihres theoretisch unendlich großen Innenwiderstands einen vom Außenwiderstand unabhängigen konstanten Strom. Neben der Realisierung hoher dynamischer Widerstände werden Konstantstromquellen auch in Schaltungen zur Pegelverschiebung eingesetzt. Prinzipiell ist eine Konstantstromquelle eine stark stromgegekoppelte Emitterschaltung, bei der im Klemmenspannungsbereich der als Laststrom auftretende Kollektorstrom fast unabhängig von der Größe des Lastwiderstands konstant bleibt.

Bei der integrierten Technik erweist sich die im Bild 3.12 angegebene Konstantstromquelle als vorteilhaft. Dabei wird die Erscheinung, daß in einem Halbleiter eng benachbarte Transistoren bei gleichen Basis-Emitter-Spannungen auch die gleiche Emitterdichte haben, ausgenutzt.

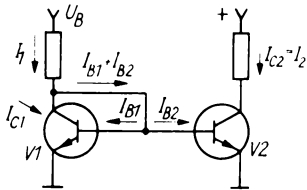


Bild 3.12.
Konstantstromquelle

Die Schaltung besteht aus zwei gleichstromgekoppten Transistoren, deren Basis-Emitter-Übergänge parallelgeschaltet sind. Basis und Kollektor des Transistors V1 sind verbunden. Nun ist das Stromverhältnis I_1/I_2 gleich dem Querschnittsverhältnis der Basis-Emitter-Strecken der Transistoren V1 und V2. I_2 ist also direkt I_1 proportional, bzw. I_1 erscheint „gespiegelt“ als I_2 im Transistor V2. Man bezeichnet Stromquellen, deren Ausgangsstrom einem anderen Strom in der Schaltung direkt proportional ist, als Stromspiegel.

Für ein Querschnittsverhältnis der Basis-Emitter-Strecken von 1 : 1 gilt folgende Überlegung: I_{B1} steuert den Kollektorstrom I_{C1} . Mit der Gleichstromverstärkung B folgt

$$I_{C1} = B I_{B1}$$

$$I_1 = I_{C1} + I_{B1} + I_{B2} = B I_{B1} + I_{B1} + I_{B2}.$$

Bei gleichen Transistorkennwerten ist $I_{B1} = I_{B2} = I_B$, damit

$$I_1 = I_B(B + 2).$$

Für $I_{C2} = I_2$ folgt analog

$$I_2 = B I_{B2} = B I_B.$$

Als Verhältnis I_2/I_1 ergibt sich:

$$\frac{I_2}{I_1} = \frac{B I_B}{I_B(B + 2)} = \frac{B}{B + 2}$$

$$I_2 = \frac{B}{B + 2} I_1 = \frac{1}{1 + \frac{2}{B}} I_1$$

Für $B \gg 1$ ist $I_2 \approx I_1$.

Durch entsprechende Wahl des Querschnittsverhältnisses bis 10 : 1 läßt sich ein entsprechendes Stromverhältnis erreichen. Damit lassen sich auch Mehrfachkonstantstromquellen, sog. *Stromquellenbänke*, aufbauen. Bild 3.13 zeigt ein Schaltungsbeispiel. Die Stromübersetzungsverhältnisse I_2/I_1 und I_3/I_1 werden durch die entsprechenden Querschnittsverhältnisse der Basis-Emitter-Strecken der Transistoren V2 und V1 bzw. V3 und V1 bestimmt.

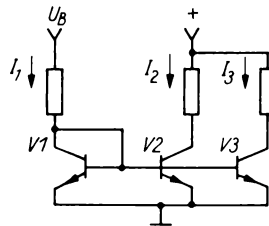


Bild 3.13.
Stromquellenbank

Referenzspannungsquellen

Referenzspannungsquellen liefern in einem bestimmten Arbeitsbereich wegen ihres sehr kleinen Innenwiderstands eine Spannung, die unabhängig vom Arbeitswiderstand, der angelegten Betriebsspannung und der Umgebungstemperatur nahezu konstant ist. Man benötigt diese in der elektronischen Schaltungstechnik sehr häufig für Vergleichs- und Meßzwecke, gelegentlich auch in Verstärkerschaltungen.

Für Referenzspannungsquellen verwendet man Bauelemente mit ausgeprägtem Knick in der Strom-Spannungs-Kennlinie (Bild 3.14). Diese

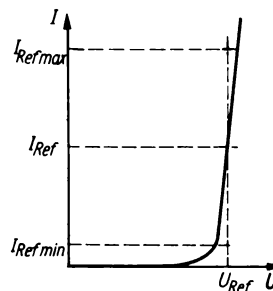


Bild 3.14. Strom-Spannungs-Kennlinie des Referenzelements

Bauelemente nennt man *Referenzelemente*. Bei ihnen ist in einem bestimmten Bereich der Spannungsabfall weitgehend unabhängig von dem sie durchfließenden Strom, so daß sich die im Bild 3.15 angegebene Grundschialtung anbietet.

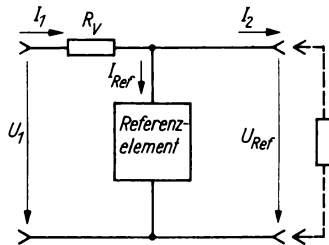


Bild 3.15. Grundschialtung einer Referenzspannungsquelle

Als Referenzelemente können im einfachsten Fall in Flußrichtung betriebene Si-Dioden oder für größere Referenzspannungen Z-Dioden (s. Abschn. 1.4.2.) verwendet werden. Bei höheren Anforderungen bezüglich des Temperaturverhaltens und eines niedrigen Ausgangswiderstands werden Transistorschaltungen eingesetzt, die sich auch problemlos in integrierter Technik herstellen lassen.

Sehr verbreitet ist eine Anordnung nach Bild 3.16. Dabei wird das Durchbruchverhalten des Emitter-Basis-Übergangs des Transistors V1 genutzt, das so wie das einer Z-Diode ist. Der Basis-Emitter-Übergang des Transistors V2 dient zur teilweisen Kompensation der Temperatureinflüsse. Außerdem wirkt V2 als Parallelregler, da der Kollektor mit dem Ausgang verbunden ist; der differentielle Widerstand r_{Ref}

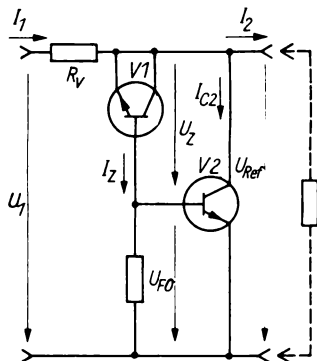


Bild 3.16. Referenzspannungsquelle mit Kompensation der Temperatureinflüsse und niedrigem Ausgangswiderstand

wird dadurch vermindert. Als Referenzspannung ergibt sich

$$U_{Ref} = U_Z + U_{F0};$$

U_{F0} Flußspannung bei einsetzendem Flußstrom (beim Si-pn-Übergang ist $U_{F0} \approx 0,7 \text{ V}$)

U_Z Durchbruchspannung des Basis-Emitter-Übergangs ($U_Z \approx 7 \text{ V}$).

Muß eine Referenzspannung mit sehr großer Temperaturstabilität bei geringer Betriebsspannung bereitgestellt werden, dann verwendet man Schaltungen, die nach dem Bandgap-Prinzip arbeiten. Bandgap ist die Breite des verbotenen Bandes eines Halbleiters nach dem Bändermodell.

Aus der Halbleiterphysik ist bekannt, daß die Basis-Emitter-Spannung eines Transistors einen negativen Temperaturkoeffizienten hat. Die Basis-Emitter-Spannungen zweier integrierter Transistoren gleicher Temperatur, die bei verschiedenen Emitterstromdichten arbeiten, haben aber einen positiven Temperaturkoeffizienten.

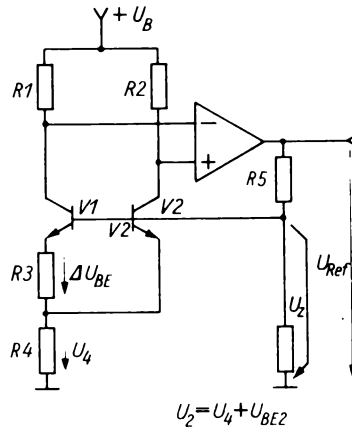


Bild 3.17. Bandgap-Referenzspannungsquelle

Bild 3.17. zeigt das Bandgap-Prinzip. Durch zwei Transistoren mit unterschiedlichen Emitterflächen fließen gleichgroße Kollektorströme. Wegen der unterschiedlichen Emitterstromdichten treten unterschiedliche Emitterstromdichten auf und es entstehen zwei verschiedene Basis-Emitter-Spannungen der Transistoren V1 und V2. Die sich ergebende Spannungsdifferenz mit positivem Temperaturkoeffizienten liegt über dem Widerstand R3 und wird mit dem Verhältnis R_4/R_3 auf R4 transformiert. Dort kompensiert sich U_{BE} mit neg. Temperaturkoeffizienten von V2 mit der auf R4 transformierten Span-

nungsdifferenz zu einer temperaturunabhängigen Summenspannung, die Bandgap-Spannung genannt wird. Die Bandgap-Spannung U_z ist die extrapolierte Spannung der Breite des verbotenen Bandes für Silizium und beträgt $U_z \approx 1,23$ V. Durch entsprechende Dimensionierung der Widerstände R_5 und R_6 wird die gewünschte Referenzspannung U_{Ref} festgelegt. Eine weitere, in integrierten Analogschaltkreisen häufig angewandte Referenzspannungsquelle, ist die Widlar-Diode. Das ist eine aus drei Transistoren und drei Widerständen bestehende Schaltung (Bild 3.18.).

Die Widlar-Diode liefert ebenfalls eine temperaturkompensierte Bandgap-Referenzspannung von $U_{Ref} \approx 1,23$ V. Das Wirkprinzip dieser Schaltung besteht in der Kompensation des Temperaturdurchgriffs des Transistors V_3 durch die Transistoren V_1 und V_2 .

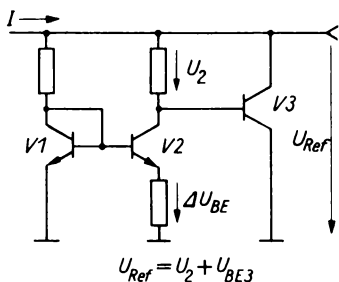


Bild 3.18.
Widlar-Diode

Koppelschaltungen

Gleichspannungsverstärker sind direktgekoppelt. Da die einzelnen Stufen voneinander abweichende Ausgangs- und Eingangsruhepotentiale haben, muß eine gleichstrommäßige Anpassung erfolgen. Bild 3.19 zeigt eine Gleichspannungskoppelschaltung.

Die aus V_1 und V_2 bestehende Konstantstromquelle verursacht an R_1 einen Gleichspan-

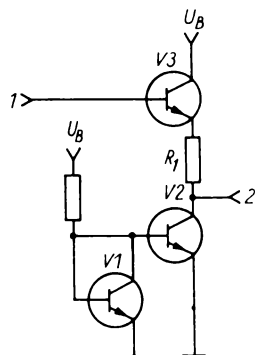


Bild 3.19. Gleichspannungskoppelschaltung

nungsabfall und verschiebt damit den Gleichspannungspegel auf den für die nachfolgende Stufe erforderlichen Wert. Ist der Eingangswiderstand der Folgestufe groß gegenüber R_1 , wird das an I angelegte Eingangssignal nicht gedämpft.

Koppelschaltungen mit Transistoren unterschiedlichen Leitungstyps ermöglichen neben der erwünschten Potentialverschiebung gleichzeitig eine Spannungsverstärkung. Bild 3.20 zeigt eine entsprechende Schaltung.

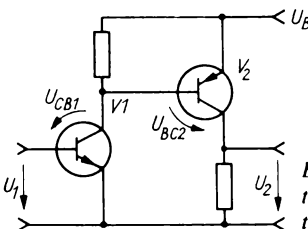


Bild 3.20. Koppelschaltung mit Komplementärtransistoren

Die beiden direktgekoppelten Komplementärtransistoren V_1 und V_2 arbeiten in Emitterschaltung. Für die Potentialverschiebung $U_v = U_2 - U_1$ erhält man $U_v = U_{CB1} - U_{BC2}$. Die Spannungsverstärkung ist $v_u = v_{u1} v_{u2}$. Wegen des einfachen Aufbaus der Schaltung, der erreichbaren Potentialverschiebung und Spannungsverstärkung werden Schaltungsstrukturen entsprechend Bild 3.20 als Koppelschaltungen in integrierten Verstärkern vorrangig eingesetzt.

Verstärkerschaltungen im Gegentaktbetrieb

Für verzerrungsarme Verstärkungen bei höheren Ausgangsleistungen und niedrigem Ruhestrom werden Gegentaktschaltungen angewendet. Bild 3.21 zeigt eine komplementäre Gegentaktschaltung, wie sie häufig als Ausgangsstufe in integrierten Schaltkreisen eingesetzt wird.

Transistor V_1 wird leitend, wenn die Eingangsspannung U_1 die Bezugsgleichspannung $U_B/2$ überschreitet; V_2 wird leitend, wenn U_1 die Spannung $U_B/2$ unterschreitet. Ist die Eingangs-

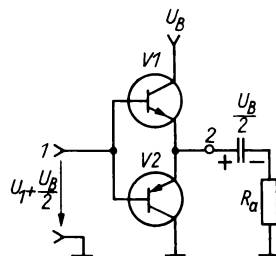


Bild 3.21. Komplementäre Gegentaktschaltung

spannung Null, bleiben beide Transistoren gesperrt.

Diese als B-Betrieb bezeichnete Betriebsart hat den Nachteil, daß besonders bei kleinen Signalen Übernahmeverzerrungen entstehen. Dieser Nachteil läßt sich durch AB-Betrieb beseitigen, indem die Ausgangstransistoren so vorgespannt werden, daß auch bei fehlendem Eingangssignal schon ein geringer Ruhestrom fließt.

Diese Schaltung erfordert ein komplementäres Transistorpaar, d. h., $V1$ und $V2$ müssen in ihrer Stromverstärkung und in ihrer Strombelastbarkeit übereinstimmen. Da das aber technologisch kompliziert ist, werden eigentlich Endtransistoren nur eines Leitfähigkeitstyps eingesetzt. Bild 3.22 zeigt eine entsprechende AB-Ausgangsstufe.

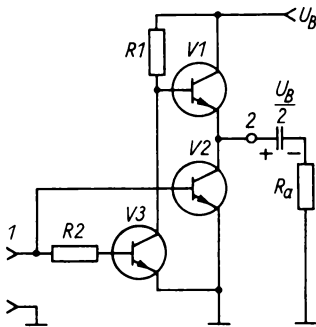


Bild 3.22. Gegentakt-AB-Ausgangsstufe

Durch das Querschnittsverhältnis der Basis-Emitter-Strecken der Transistoren $V3$ und $V2$ ist der Ruhestrom bestimmt. Da der Spannungsabfall an $R2$ im Ruhezustand vernachlässigbar klein bleibt, erscheint der durch $V3$ fließende Strom auch in $V2$ gespiegelt. Durch den Widerstand $R1$ und den durch den Transistor $V3$ fließenden Strom wird das Ausgangsruhepotential auf $U_B/2$ eingestellt. Wird die Eingangsspannung U_1 positiv, dann wird der Transistor $V2$ leitend, und die Ausgangsspannung sinkt unter den Bezugswert von $U_B/2$; wird dagegen U_1 negativ, dann werden $V3$ und $V2$ zunehmend gesperrt, und die Ausgangsspannung steigt über den Bezugswert von $U_B/2$. Wenn bei einer großen negativen Ansteuerung $V3$ und $V2$ völlig gesperrt werden, wird über $R1$ der Transistor $V1$ voll durchgesteuert und liefert Strom in die Last. Übernahmeverzerrungen treten bei dieser Schaltung nicht auf, da die Ausgangstransistoren auch im Ruhezustand Strom führen.

Zusammenfassung

Ein Verstärker ist ein aktiver Vierpol mit zwei Eingangs- und zwei Ausgangsbuchsen.

Beim *Vorverstärker* strebt man einen sehr hohen Verstärkungsgrad an, während der Wirkungsgrad wegen der kleinen Leistungen uninteressant und meist gering ist. Auf Grund der kleinen Steueramplitude arbeiten Vorverstärker in erster Näherung verzerrungsfrei.

Beim *Endverstärker* strebt man die höchste verzerrungsfreie Ausgangsleistung an. Hoher Verstärkungs- und Wirkungsgrad sind als zweitrangig zu betrachten.

Beim *Senderverstärker* muß auf einen hohen Wirkungsgrad größter Wert gelegt werden. Weniger wichtig sind Verzerrungsfreiheit und große Verstärkung.

Zur Beurteilung von Verstärkern dienen mehrere Kenngrößen. Die wichtigsten sind Verstärkungsgrad v , Frequenzbandbreite b und Klirrfaktor k . Der Verstärkungsgrad gibt stets das Verhältnis von Ausgangs- zu Eingangsgröße an. Die Frequenzbandbreite gibt an, welcher Frequenzbereich annähernd gleich groß verstärkt wird. Die Eckfrequenzen nennt man untere bzw. obere Grenzfrequenz.

Der Klirrfaktor charakterisiert die nichtlinearen Verzerrungen und ergibt sich aus dem Verhältnis des Effektivwerts aller Oberwellen zu dem aller Harmonischen.

Die Verstärkergrundsaltungen mit einem Transistor werden nach der Ein- und Ausgangsgemeinsamen Elektrode des Verstärkerbauelements benannt. Alle Grundsaltungen haben voneinander abweichende Eigenschaften; sie sind in Tafel 3.2. zusammengestellt.

Die Bootstrap-Schaltung arbeitet nach einem Schaltungsprinzip, bei dem die Spannungen zweier Punkte „mitlaufen“. Damit kann u. a. ein Widerstand dynamisch vergrößert werden. Erweiterte Grundsaltungen benötigen zwei und mehr Transistoren. Mit ihnen lassen sich speziellere Forderungen erfüllen.

Die Kaskaden- oder Darlington-Schaltung kann als Standard- oder Komplementärschaltung realisiert werden. Mit ihr ist eine große Stromverstärkung und ein hoher Eingangswiderstand erreichbar.

Die Kaskodeschaltung hat eine wesentlich geringere dynamische Eingangskapazität als eine einfache Emitterschaltung und kann deshalb vorteilhaft als Eingangsstufe in Hochfrequenzverstärkern eingesetzt werden.

Die Kombination Unipolar- und Bipolartransistor ermöglicht einen sehr hohen Eingangswiderstand und ein günstiges Rauschverhalten eines Verstärkers.

Die Differenzverstärkerschaltung ist eine Grundstruktur, die sehr häufig insbesondere bei integrierten Analogschaltkreisen angewendet wird. Die Schaltung mit zwei Ein- und zwei Ausgängen verstärkt gegenphasig Eingangssignale hoch, dagegen gleichphasig nur sehr gering. Demzufolge unterscheidet man zwischen einer Differenz- und einer Gleichtaktverstärkung. Wegen der geringen Gleichtaktverstärkung ist die Differenzverstärkerschaltung bei Temperaturänderungen wesentlich weniger empfindlich als alle anderen Verstärkerschaltungen.

Referenzspannungsquellen haben in einem bestimmten Arbeitsbereich einen sehr kleinen Innenwiderstand. Sie liefern eine weitgehend konstante Vergleichsspannung.

Konstantstromquellen sollen einen von der Belastung unabhängigen Strom liefern.

Koppelschaltungen werden in direktgekoppelten Verstärkern zur Potentialverschiebung benötigt. Besonders vorteilhaft lassen sich derartige Schaltungen mit Komplementärtransistoren aufbauen, da dabei neben der erwünschten Potentialverschiebung auch eine Spannungsverstärkung erfolgt.

Verstärkerschaltungen im Gegentaktbetrieb ermöglichen bei niedrigem Ruhestrom verzerrungsarme Verstärkungen und höhere Ausgangsleistungen. Ausgangsstufen integrierter Schaltkreise sind sehr häufig als Gegentakverstärker aufgebaut.

3.3.

Gegenkopplung

3.3.1.

Allgemeines

Wird ein Teil der Ausgangsgröße eines Verstärkers auf dessen Eingang zurückgeführt, dann spricht man allgemein von einer Rückkopplung. Ist diese rückgekoppelte Größe dem Eingangssignal entgegengerichtet und tritt eine Verstärkungsminderung auf, handelt es sich um eine *Gegenkopplung* (negative Rückkopplung), im umgekehrten Fall um eine Mitkopplung.

Gegenkopplungsmaßnahmen werden ihrer Vorteile wegen in Verstärkern außerordentlich oft

angewendet. Die wichtigsten Eigenschaften einer Gegenkopplung sind:

- Durch Gegenkopplung werden die nichtlinearen Verzerrungen eines Verstärkers vermindert, da die Steuerkennlinie des aktiven Verstärkerbauelements linearisiert wird. Die neue Kennlinie hat gegenüber der ursprünglichen einen längeren „geradlinigen“ Teil; sie knickt aber auch schärfer ab. Das ist bei Endverstärkern besonders zu beachten, da sonst bei schon geringfügiger Übersteuerung der Klirrfaktor stark anwächst.
- Gegengekoppelte Verstärker weisen gegenüber nichtgegengekoppelten geringere Verstärkungsschwankungen auf, die durch Altern der Bauelemente, Einsatz von Bauelementen mit unterschiedlichen Verstärkungsfaktoren, Betriebsspannungsschwankungen und Temperaturänderungen hervorgerufen werden.
- Durch Gegenkopplung lassen sich in starkem Maße die Ausgangs- und Eingangswiderstände der Verstärker verändern.
- Bedingt durch die obengenannte Eigenschaft kann man einerseits bei gegengekoppelten Verstärkern eine größere Frequenzbandbreite und einen gleichmäßigeren Frequenzgang als bei nichtgegengekoppelten erreichen. Durch eine frequenzabhängige Gegenkopplung kann man andererseits einen beliebigen Frequenzverlauf erzielen, z. B. eine stärkere Schwächung der mittleren Frequenzen gegenüber den tiefen und hohen.
- Nachteilig ist bei einem gegengekoppelten Verstärker der geringere Verstärkungsfaktor. Benötigt man die gleiche Ausgangsspannung oder -leistung, ist also ein größeres Eingangssignal erforderlich.

3.3.2.

Ersatzschaltungen des Verstärkerbauelements

Das Wechselstromverhalten eines Verstärkerbauelements läßt sich relativ einfach beschreiben, wenn man es als einen aktiven Vierpol betrachtet (s. Bild 3. 1). Man untersucht dabei lediglich, wie es sich bezüglich der beiden Eingangs- und Ausgangsgrößen an den Anschlußpunkten verhält, ohne zu berücksichtigen, was im einzelnen im Innern des Vierpols vorgeht und wie er aufgebaut ist. Allerdings müssen lineare Verhältnisse vorausgesetzt werden, die

aber bei Vorverstärkern wegen der kleinen Aussteuerung mit guter Näherung zutreffen.

Zwei häufig benutzte Ersatzschaltungen sind die y - und die h -Ersatzschaltung.

Bei Transistoren in HF-Verstärkern benutzt man meist die y -Ersatzschaltung, bei Transistoren im NF-Gebiet am häufigsten die h -Ersatzschaltung.

Zu beachten ist dabei, daß die Parameter nur in niedrigem und mittlerem Frequenzbereich reell sind. Die sich anschließenden Ausführungen gelten nur unter der einschränkenden Voraussetzung reeller Parameter.

Die Betrachtungen beziehen sich auf die h -Ersatzschaltung (Tafel 3.3.). Für den Eingangs-

und Ausgangskreis gelten folgende Beziehungen:

$$u_1 = U_0 - i_1 R_0 \quad (3.17)$$

U_0 ~ Leerlaufspannung der Steuerspannungsquelle
 R_0 ~ Innenwiderstand der Steuerspannungsquelle.

$$u_2 = -i_2 R_a \quad (3.18)$$

Zu beachten ist dabei, daß die Beziehungen allgemeingültig sind und sich auf keine bestimmte Grundschialtung beschränken. Für die betreffenden Schaltungen sind lediglich die für sie gültigen h -Parameter zu verwenden. Sie werden durch Indizes benannt, und zwar für die

Emitterschaltung mit dem Index e
 Basisschaltung mit dem Index b
 Kollektorschaltung mit dem Index c.

Durch Gegenkopplung können die h -Parameter des Verstärkerbauelements stark verändert werden. Es ergeben sich damit auch andere Betriebsgrößen als bei einem nichtgegekoppelten Verstärker. Durch geeignete Dimensionierung des Gegenkopplungsnetzwerks lassen sich damit Verstärkereigenschaften erreichen, die den praktischen Erfordernissen entsprechen.

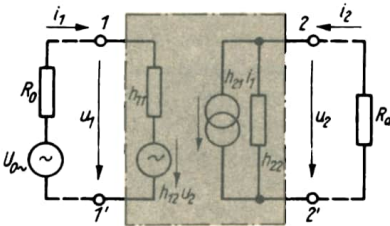
Die in Tafel 3.4 angegebenen Berechnungsgrundlagen gelten auch für gegengekoppelte Verstärker, wenn die dafür gültigen h' -Parameter eingesetzt werden. Im Abschn. 3.3.3. erfolgt eine genauere Begründung.

Tafel 3.4. Betriebsgrößen eines Transistorvierpols (h -Ersatzschaltung)

Die Gleichungen sind Näherungen für $h_{12} = 0$

Ströme	$i_1 \approx U_0 \sim \frac{1}{h_{11} + R_0}$
	$i_2 \approx U_0 \sim \frac{h_{21}}{(h_{11} + R_0)(1 + h_{22}R_a)}$
Spannungen	$u_1 \approx U_0 - \left(1 - \frac{R_0}{h_{11} + R_0}\right)$
	$u_2 \approx -U_0 \sim \frac{h_{21}R_a}{(h_{11} + R_0)(1 + h_{22}R_a)}$
Widerstände	$R_{\text{ein}} \approx h_{11}$
	$R_{\text{aus}} \approx \frac{1}{h_{22}}$
Verstärkungen	$v_i = \frac{h_{21}}{1 + h_{22}R_a}$
	$v_u \approx \frac{-h_{21}R_a}{h_{11}(1 + h_{22}R_a)}$
	$v_p \approx \frac{h_{21}^2 R_a}{h_{11}(1 + h_{22}R_a)^2}$

Tafel 3.3. Die h -Ersatzschaltung des Transistors



Dieser Vierpol wird durch die beiden Gleichungen

$$u_1 = h_{11}i_1 + h_{12}u_2 \quad (3.13)$$

$$i_2 = h_{21}i_1 + h_{22}u_2 \quad (3.14)$$

beschrieben.

Die h -Parameter (h = hybrid, gemischt) stellen sowohl Widerstände als auch Leitwerte und einheitenlose Kenngrößen dar.

h_{11} ist der Eingangswiderstand bei kurzgeschlossenem Ausgang ($u_2 = 0$)

$$h_{11} = \frac{u_1}{i_1} \bigg|_{u_2 = 0} \quad (3.15.a)$$

h_{22} ist der Ausgangsleitwert bei offenem Eingang ($i_1 = 0$)

$$h_{22} = \frac{i_2}{u_2} \bigg|_{i_1 = 0} \quad (3.16.a)$$

h_{12} ist die Spannungsrückwirkung bei offenem Eingang ($i_1 = 0$)

$$h_{12} = \frac{u_1}{u_2} \bigg|_{i_1 = 0} \quad (3.16.b)$$

h_{21} ist die Stromverstärkung bei kurzgeschlossenem Ausgang ($u_2 = 0$)

$$h_{21} = \frac{i_2}{i_1} \bigg|_{u_2 = 0} \quad (3.15.b)$$

3.3.3.

Gegenkopplungsschaltungen

Im Bild 3.23 wird oben ein gegengekoppelter Verstärker mit der Verstärkung v' wiedergegeben. Die Eingangsgrößen sind U'_{1-} und I'_{1-} , die Ausgangsgrößen U'_{2-} und I'_{2-} . Dieser Verstärkervierpol besteht aus einem nichtgegengekoppelten Verstärker mit der Verstärkung v und einem Gegenkopplungsvierpol, der in vier Varianten mit ersteren zusammenschlossen werden kann. In allen Fällen wird auf den Verstärkereingang der k -te Teil der Ausgangsgröße geführt. Ist die zurückgeführte Größe der Ausgangsspannung proportional, dann liegt eine *Spannungsgegenkopplung* vor. Das Gegenkopplungsnetzwerk liegt dabei parallel zum Ausgang. Ist dagegen die zurückgeführte Größe dem Ausgangsstrom proportional, dann ist das eine *Stromgegenkopplung*. Das Gegenkopplungsnetzwerk liegt dabei in Serie (Reihe) zum Ausgang.

Der Gegenkopplungsvierpol kann sowohl eine als auch mehrere Verstärkerstufen überbrücken. Wichtig ist nur, daß die zurückgeführte Größe dem Eingangssignal entgegenwirkt. Auch hier kann durch den Einsatz von Blindschaltetelementen eine frequenzabhängige Gegenkopplung aufgebaut und ein ungleichmäßiger Frequenzgang erreicht werden.

Die Gegenkopplungsschaltungen werden durch zwei Wörter gekennzeichnet, wobei sich das erste auf den Eingang und das zweite auf den Ausgang bezieht. Man erhält dann die

- Serien-Parallel-Gegenkopplung (a)
- Parallel-Serien-Gegenkopplung (b)
- Serien-Serien-Gegenkopplung (c)
- Parallel-Parallel-Gegenkopplung (d).

Hauptbedeutung haben in der Verstärkerpraxis die Varianten c und d.

Die Rückkopplungsgrundgleichung (3.19) läßt sich leicht aus Bild 3.23 a ableiten

$$v' = v'_u = \frac{U_{2-}}{U'_{1-}} \quad (3.19 \text{ a})$$

$$U'_{1-} = U_{1-} - kU_{2-}. \quad (3.19 \text{ b})$$

Gl. (3.19 b) wird in Gl. (3.19 a) eingesetzt und mit $1/U_{1-}$ erweitert:

$$v' = v'_u = \frac{U_{2-}}{U_{1-} - kU_{2-}} = \frac{\frac{U_{2-}}{U_{1-}}}{1 - k \frac{U_{2-}}{U_{1-}}}.$$

Da $v = v_u = U_{2-}/U_{1-}$ die Verstärkung ohne Gegenkopplung ist, folgt sofort

$$v = \frac{v}{1 - kv} \quad (3.19)$$

Zum gleichen Ergebnis kommt man, wenn man die Parallel-Serien-Gegenkopplung nach Bild 3.23 b zugrunde legt, nur hat man dort für $v' = v'_i = I_{2-}/I'_{1-}$ und für $v = v_i = I_{2-}/I_{1-}$ einzusetzen.

In dieser wichtigen Grundgleichung (3.19) wird das Produkt kv als *Schleifenverstärkung* bezeichnet.

Der Ausdruck $1 - kv = v/v'$ heißt *Rückkopplungsgrad*. Für den Fall, daß der Betrag des Rückkopplungsgrads

- | $1 - kv$ | > 1 ist, liegt eine Gegenkopplung vor, dagegen
- für | $1 - kv$ | < 1 eine Mitkopplung, die bei
- | $1 - kv$ | = 0 zur Selbsterregung der Schaltung führt.

Die Gegenkopplung verändert aber auch den Ausgangs- und Eingangswiderstand eines Verstärkers, und zwar um so mehr, je größer der Rückkopplungsgrad ist. Bei einer Stromgegenkopplung nach Bild 3.23 c beispielsweise würde bei abgetrenntem Arbeitswiderstand kein Ausgangsstrom mehr fließen können. Damit würde die am Gegenkopplungswiderstand R_I abfallende Gegenkopplungsspannung verschwinden, die innere Verstärkung steigen und mit ihr bei konstantem Eingangssignal die Ausgangsspannung. Dieses Verhalten zeigen aber Spannungsquellen mit großem Innenwiderstand, und man kann verallgemeinernd feststellen, daß eine Serien- oder Stromgegenkopplung den Ausgangswiderstand eines Verstärkers vergrößert.

Der Eingangswiderstand der gleichen Schaltung wird auch mit zunehmendem Rückkopplungsgrad vergrößert; denn der Widerstand zwischen Basis und Emitter liegt nicht allein, sondern in Reihe mit R_I an der Eingangsspannung U'_{1-} und erscheint dadurch vergrößert.

Ein spannungsgegengekoppelter Verstärker zeigt bezüglich des Ausgangs- und Eingangswiderstands ein anderes Verhalten. Beispielsweise soll in der Schaltung nach Bild 3.23 d die Belastung verringert werden, indem der Arbeitswiderstand R_a vergrößert oder ganz abgetrennt wird. Damit steigt auch die Ausgangsspannung. Letztere ist aber der Gegenkopplungsspannung proportional, so daß auch diese ansteigt und die

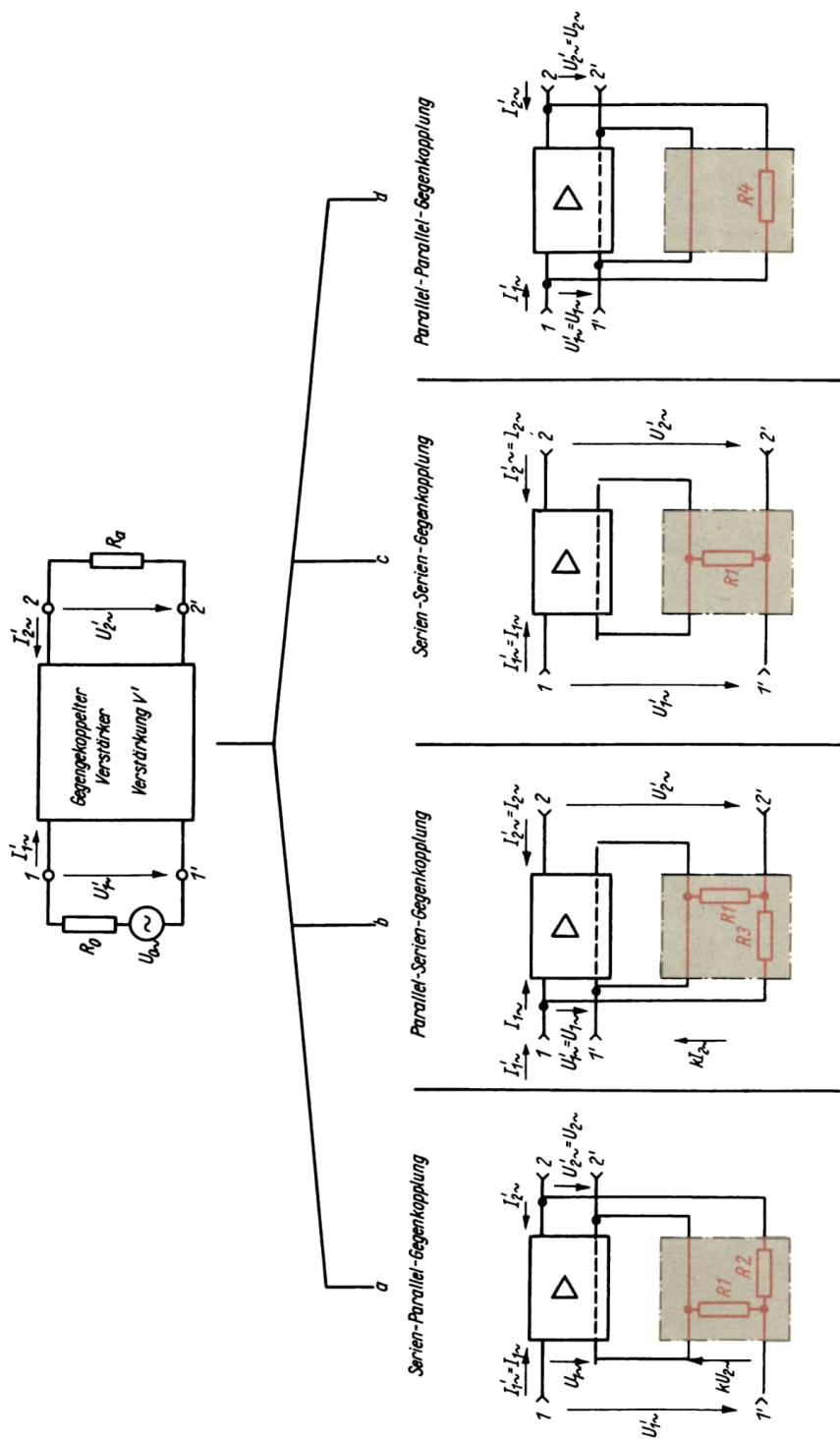


Bild 3.23. Gegenkopplungsgrundschaltungen

Verstärkung mehr als vorher herabsetzt. Die Folge davon ist ein nur geringer Ausgangsspannungsanstieg.

Spannungsquellen mit einem kleinen Innenwiderstand zeigen eine ebensolche geringe Belastungsabhängigkeit der Ausgangsspannung, und man kann feststellen, daß eine Parallel- oder Spannungsgegenkopplung den Ausgangswiderstand des Verstärkers verkleinert.

Der Eingangswiderstand der gleichen Schaltung wird wegen der Stromteilung durch den Gegenkopplungswiderstand im Eingangskreis auch vermindert. Auf Grund der Verstärkung liegt über R_4 eine beträchtliche Spannung. Diese treibt einen großen Strom an und kann zu einer starken Belastung der Steuerquelle führen.

Im Bild 3.24 wird das prinzipielle Verhalten gegengekoppelter Verstärker wiedergegeben und die Spannungsverstärkung, Stromverstärkung, der Ausgangs- und Eingangswiderstand als Funktion des Rückkopplungsgrads der vier Grundschaltungen dargestellt.

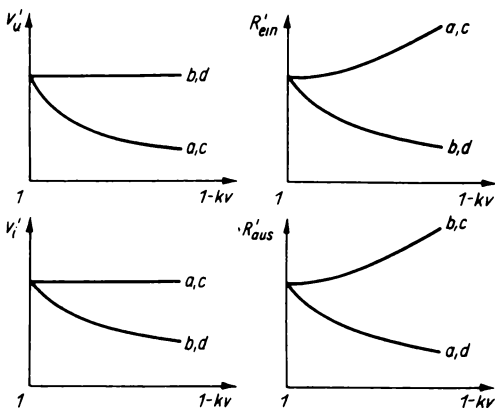


Bild 3.24. Verhalten gegengekoppelter Verstärker in Abhängigkeit vom Rückkopplungsgrad

- a) Serien-Parallel-Gegenkopplung
- b) Parallel-Serien-Gegenkopplung
- c) Serien-Serien-Gegenkopplung
- d) Parallel-Parallel-Gegenkopplung

Zusammenfassung

Gegenkopplung ist eine negative Rückkopplung. Bei ihr wird ein Teil der Ausgangsgröße eines Verstärkers so auf dessen Eingang zurückgeführt, daß diese Größe dem Eingangssignal entgegenwirkt.

Neben dem Nachteil eines geringeren Verstärkungsfaktors haben gegengekoppelte gegenüber nichtgegengekoppelten Verstärkern die Vor-

teile, daß sie kleinere nichtlineare Verzerrungen aufweisen, weniger empfindlich auf Bauelementalterung, Betriebsspannungsschwankungen und Temperaturänderungen reagieren und eine größere Frequenzbandbreite bei gleichmäßigem Frequenzgang haben.

Außerdem läßt sich durch eine frequenzabhängige Gegenkopplung ein beliebiger Frequenzverlauf erzielen. Die Ausgangs- und Eingangswiderstände eines Verstärkers lassen sich durch Gegenkopplung in starkem Maße variieren. Ist die rückgekoppelte Größe der Ausgangsspannung proportional, dann liegt eine Spannungsgegenkopplung vor. Ist sie dagegen dem Ausgangsstrom proportional, dann spricht man von einer Stromgegenkopplung.

Es gibt vier Grundschaltungen, die nach der Art der Anschaltung des Gegenkopplungsvierpols an den Eingangs- und Ausgangskreis gekennzeichnet werden. Sie lassen sich mathematisch durch die Gleichung

$$v' = \frac{v}{1 - kv}$$

beschreiben. Eine Gegenkopplung liegt dann vor, wenn $|1 - kv| > 1$ ist; anderenfalls handelt es sich um eine Mitkopplung.

Hinsichtlich der Eingangs- und Ausgangswiderstände gilt der allgemeine Sachverhalt, daß sie bei einer Seriengegenkopplung vergrößert, bei einer Parallelgegenkopplung dagegen verkleinert werden.

3.4. Operationsverstärker

3.4.1. Allgemeines

Operationsverstärker (OPV) wurden für die Durchführung mathematischer Operationen in der Analogrechenstechnik entwickelt und werden deshalb auch als Rechenverstärker bezeichnet. Es sind standardisierte Universalschaltungen, um gewünschte Eingangs- und Ausgangsbeziehungen zu realisieren. Operationsverstärker werden vor allem als integrierte Schaltkreise ausgeführt, wobei der bipolare Funktionsblock dominiert. Es handelt sich dabei um mehrstufige Gleichstromverstärker. Sie werden als ein Bauelement behandelt und mit starker äußerer Gegenkopplung betrieben. Durch die zugeschalteten diskreten Bauelemente kann man sehr un-

terschiedliche Eigenschaften der Gesamtanordnung erreichen.

Allen OPV ist eine sehr hohe Verstärkung, große Bandbreite und ein charakteristischer Frequenzgang gemeinsam. Technisch werden Operationsverstärker in drei Arten ausgeführt:

- OPV mit nur einem stets invertierend wirkenden Eingang. Die Ausgangsspannung ist dabei zur Eingangsspannung gegenphasig.
- OPV mit zwei Eingängen, wobei der eine invertierend, der andere nichtinvertierend wirkt.

Verstärkt wird nur die zwischen beiden Eingängen liegende Differenzspannung.

- OPV mit Gegentaktausgang. Bei diesem relativ selten angewendeten OPV liefert eine zwischen den Eingangsklemmen liegende Differenzspannung zwei entsprechend verstärkte Ausgangsspannungen entgegengesetzter Polarität.

Alle Eingangs- und Ausgangsspannungen werden auf Massepotential bezogen. In zeichnerischen Darstellungen wird ein OPV durch ein gleichseitiges Dreieck wiedergegeben; die Eingänge werden an der vertikalen Seite gekennzeichnet (Bild 3.25).

Da das Ausgangssignal des OPV sowohl positive als auch negative Werte annehmen kann, sind auch zwei entgegengerichtete Betriebsspannungen nötig. Diese haben in der Regel den gleichen Absolutwert.

Um eine unerwünschte Selbsterregung über die Betriebsspannungszuführung zu vermeiden, sind unmittelbar an den Betriebsspannungsanschlüssen eines integrierten Operationsverstärkers HF-Kondensatoren (z. B. keramische Kondensatoren) mit einer Kapazität von 10 bis 100 nF nach Masse zu schalten (Bild 3.26).

Neben der im Bild 3.26 angegebenen Schaltungsmaßnahme gibt es zahlreiche Schutzschaltungen für den Schaltkreis. Einige sind im Bild 3.27 angegeben. Sie können im konkreten Fall kombiniert oder modifiziert werden. Die Schaltung nach Bild 3.27 a schützt den Eingang eines Spannungsfolgers gegen Überspannungen. Bild 3.27 b zeigt eine Schutzschaltung gegen unzulässig hohe Differenzeingangsspannungen für invertierende Verstärker. Ein Widerstand am Ausgang entspr. Bild 3.27 c schützt den Schaltkreis vor zu hohen Lastströmen. Bild 3.27 d zeigt eine Schaltung, die den Ausgang vor Überspannungen schützt, die von der Last eingekoppelt werden können.

Dem Einsatz des OPV sind aber auch Grenzen gesetzt. Diese sind technisch, vor allem aber ökonomisch bedingt. Ein universeller Operationsverstärker weist nämlich viele Eigenschaften auf, die bei vielen Anwendungsfällen nur zum Teil genutzt werden. Oft genügt ein einfacher, aber spezieller Schaltkreis den gestellten Anforderungen. Ein solcher Schaltkreis ist billiger, wenn er in großen Stückzahlen produziert und eingesetzt werden kann.

Eine weitere Schwierigkeit ist, daß OPV bei höheren Frequenzen nur bedingt verwendet werden können. OPV-Schaltungen sind stark gekoppelt. Damit die Gegenkopplung nicht

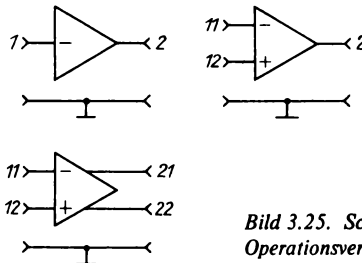


Bild 3.25. Schaltzeichen für Operationsverstärker

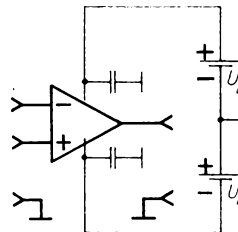


Bild 3.26. Stromversorgung des Operationsverstärkers

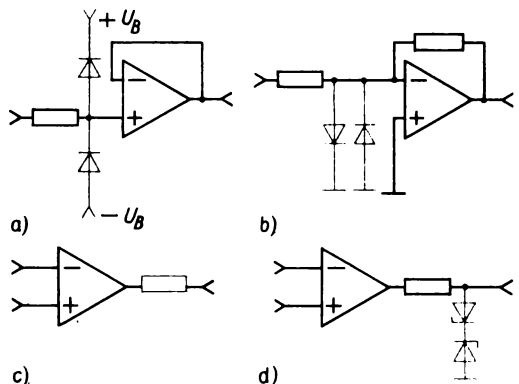


Bild 3.27. Schutzschaltungen für Operationsverstärker

a) gegen Überspannungen; b) gegen zu hohe Differenzeingangsspannungen; c) gegen Überlastung; d) gegen Überspannungen am Ausgang

in eine Mitkopplung umschlägt, darf der bei höheren Frequenzen einsetzende Verstärkungsabfall bis zur Verstärkung $v = 1$ nicht steiler als proportional zu $1/f$ erfolgen, also 20 dB/Dekade. Die bei $v = 1$ vorliegende Frequenz heißt Transitfrequenz f_T . Soll nun beispielsweise durch äußere Beschaltung ein OPV mit $f_T = 1$ MHz auf eine Verstärkung von $v = 100$ gebracht werden, dann tritt wegen des $1/f$ -Frequenzganges diese Verstärkung erst bei Frequenzen < 10 kHz auf.

Schlußfolgernd ergibt sich damit, daß Schaltungen mit Operationsverstärkern dann besonders vorteilhaft sind, wenn bei kurzen Entwicklungszeiten Baugruppen für nicht zu hohe Frequenzen bei kleinen bis mittleren Stückzahlen zu fertigen sind.

3.4.2.

Idealer Operationsverstärker

Die Eingangsspannung u_{IN} liegt am invertierenden, u_{IP} am nichtinvertierenden Eingang des OPV. Zwischen den beiden Eingangsklemmen wirkt damit die Differenzspannung $u_D = u_{IP} - u_{IN}$ und tritt, um den Faktor v_0 verstärkt, als $u_2 = v_0 (u_{IP} - u_{IN})$ am Ausgang auf. Die zwischen Ausgangs- und Eingangsgröße herrschende Beziehung gibt die Übertragungsfunktion $u_2 = f(u_{IP} - u_{IN})$ wieder.

Der ideale OPV weist folgende Eigenschaften auf:

- Die Spannungsverstärkung in offener Schleife (offene Spannungsverstärkung) v_0 ist unendlich.
- Die Bandbreite reicht von Null bis unendlich, d. h., die Verstärkung bleibt von der Frequenz Null bis zu beliebig hohen Frequenzen konstant.

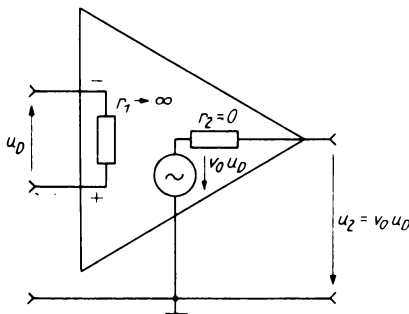


Bild 3.28. Ersatzschaltung des idealen Operationsverstärkers

- Der Eingangswiderstand r_1 ist unendlich, so daß keine Steuerleistung aufgebracht werden muß.
- Der Ausgangswiderstand r_2 ist Null, so daß sich bei Änderung des Lastwiderstands die Ausgangsspannung nicht ändert.
- Es gibt keine Nullpunktfehler, und die Eigenschaften des idealen OPV sind von Drift unabhängig.

Die Ersatzschaltung des idealen OPV wird im Bild 3.28 gezeigt.

3.4.3.

Realer Operationsverstärker

Der reale OPV kann die Eigenschaften des idealen OPV nur annähernd erreichen. Für integrierte bipolare Anordnungen ergeben sich je nach OPV-Typ folgende Kennwerte:

Offene

Spannungsverstärkung	v_0	$= 10^4 - 10^5$
Eingangswiderstand	r_1	$= 0,2 - 3 \text{ M}\Omega$
Ausgangswiderstand	r_2	$= 75 - 500 \Omega$
Eingangsfehlspannung (Offsetspannung)	U_{off}	$= 1 - 15 \text{ mV}$

Mit diesen Werten und der Schaltung eines gegengekoppelten OPV im invertierenden Betrieb nach Bild 3.29 bei einem Ausgangsspannungsbereich $U_{2\text{max}} = \pm 10 \text{ V}$ ergeben sich folgende Verhältnisse:

$$\bullet \quad |u_{IN} - u_{IP}| = \frac{|U_{2\text{max}}|}{v_0} < \frac{10 \text{ V}}{10^4} = 1 \text{ mV}$$

Damit wird die Differenzspannung u_D bei dieser großen Verstärkung vernachlässigbar gering, d. h., solange keine Übersteuerung vorliegt, fließt über r_1 kaum ein Strom, r_1 ist daher unwirksam, also

$$r_1 \rightarrow \infty$$

- Da wegen $|u_{IN} - u_{IP}| \approx 0$ auch der Eingangsstrom i_1 vernachlässigbar klein ist, erhält der invertierende Eingang scheinbar Massepotential (virtuelle Masse). Der durch R_1 fließende Strom I_1 muß entgegengesetzt gleich dem durch R_2 fließenden Strom I_2 sein, also

$$\frac{U_1}{R_1} \approx - \frac{U_2}{R_2}$$

$$\text{bzw.} \quad v = \frac{U_2}{U_1} = - \frac{R_2}{R_1}$$

Die Verstärkung v ist also nur vom Widerstandsverhältnis R_2/R_1 abhängig.

- Für den Eingangswiderstand R_{ein} folgt

$$R_{\text{ein}} = \frac{U_1}{I_1} = \frac{U_1}{U_1/R_1} = R_1$$

- Der Ausgangswiderstand R_{aus} wird vom Rückkopplungsgrad v_0/v bestimmt. Da es sich nach Bild 3.29 um eine Spannungsgegenkopplung handelt, wird bei Belastungsänderung die Ausgangsspannung nur gering verändert. Quellen mit einem kleinen Innenwiderstand zeigen eine ebensolche geringe Belastungsabhängigkeit der Ausgangsspannung. Als Näherung ergibt sich

$$R_{\text{aus}} \approx \frac{r_2}{v_0/v}$$

Bei einem Rückkopplungsgrad $v_0/v = 10^2 \dots 10^4$ und einem mittleren Ausgangswiderstand $r_2 = 100 \Omega$ wird

$$R_{\text{aus}} \approx \frac{100 \Omega}{10^2 \dots 10^4} = 10 \text{ m} \Omega \dots 1 \Omega,$$

also vernachlässigbar klein.

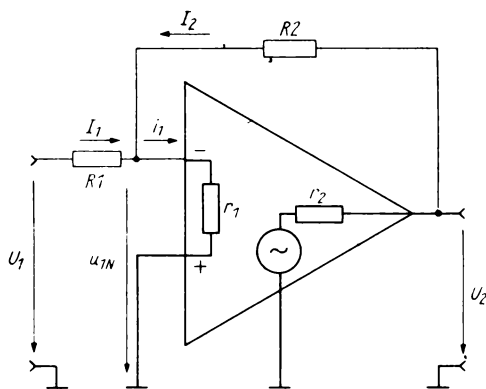


Bild 3.29. Invertierender Operationsverstärker mit Gegenkopplung

Es gibt beim realen OPV weitere Abweichungen vom idealen OPV, die aber schaltungstechnisch kompensiert werden können. Sie sollen hier nur kurz angedeutet werden; eine genauere Betrachtung erfolgt in den Abschnitten 3.4.5. und 3.4.6.

- Frequenzgang

Wie bei jedem Verstärker nimmt auch beim OPV mit zunehmender Frequenz die offene Spannungsverstärkung v_0 ab, die mit einer zusätzlichen Phasendrehung verbunden ist. Da ein Zusatzphasenwinkel $\varphi = -180^\circ$ die

Gegenkopplung zu einer Mitkopplung führt, muß durch eine äußere Beschaltung dafür gesorgt werden, daß die Schwingbedingung nicht erfüllt wird.

- Offset

Durch geringe Unsymmetrien im OPV und unterschiedliche Spannungsabfälle, die durch die Eingangsleichströme an den vor den Eingängen zugeschalteten Widerständen auftreten, ist die Ausgangsspannung $u_2 \neq 0$, wenn die Differenzspannung $u_D = u_{1P} - u_{1N} = 0$. Diese Erscheinung nennt man *Offset*. Der Offset muß fast immer beachtet und durch Gegenschaltung einer kleinen Spannung kompensiert werden.

- Gleichtaktverhalten

Werden die an beiden Eingängen liegenden Spannungen u_{1N} und u_{1P} gemeinsam geändert, also in Gleichtakt angesteuert, so bleibt beim idealen OPV die Symmetrie erhalten und die Ausgangsspannung ist $u_2 = 0$. Beim realen OPV tritt aber eine kleine Störspannung auf, die sich vor allem beim nichtinvertierenden Verstärker nachteilig auswirkt.

Schlußfolgernd ergibt sich aus den obenstehenden Betrachtungen, daß ein realer OPV in Näherung wie ein idealer OPV behandelt werden kann. Voraussetzung ist, daß er nicht übersteuert, Frequenzgang und Offset kompensiert und bei nichtinvertierendem Betrieb ein OPV mit gutem Gleichtaktverhalten ausgewählt wird.

3.4.4.

Grundsaltungen mit Operationsverstärkern

Invertierender Verstärker

Der invertierende Verstärker (Bild 3.30) wird in der Elektronik eingesetzt, da mit ihm zahlreiche Schaltungsprobleme gelöst werden können.

Unter der Annahme $v_0 \rightarrow \infty$

$$R_2 \gg r_2$$

$$r_1 \gg R_1 \parallel (R_2 + r_2)$$

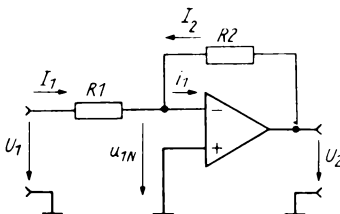


Bild 3.30. Invertierender Verstärker

gilt $i_1 = 0$ und $u_D = 0$; dann ist

$$I_1 + I_2 = 0 \text{ bzw. } I_1 = -I_2,$$

mit $I_1 = \frac{U_1}{R_1}$ und $I_2 = \frac{U_2}{R_2}$ folgt

$$\frac{U_1}{R_1} = -\frac{U_2}{R_2} \quad U_2 = -\frac{R_2}{R_1} U_1$$

$$v = \frac{U_2}{U_1} = -\frac{R_2}{R_1}$$

(3.20)

$$R_{\text{ein}} = \frac{U_1}{I_1} \text{ und } I_1 = \frac{U_1}{R_1}$$

$$R_{\text{ein}} = \frac{U_1}{U_1/R_1} = R_1$$

(3.21)

Für den Ausgangswiderstand gilt allgemein:

$$R_{\text{aus}} = \frac{U_{21}}{I_{2k}};$$

U_{21} Leerlaufspannung ($R_s \rightarrow \infty$)

I_{2k} Kurzschlußstrom ($R_s = 0$).

$$U_{21} = -\frac{R_2}{R_1} U_1$$

$$I_{2k} = \frac{-v_0 u_D}{r_2}$$

da $U_2 = 0$ (Kurzschluß), ist $I_2 = 0$, d. h.,

$$u_D = u_{1N} = U_1 \frac{R_2}{R_1 + R_2} \text{ und}$$

$$I_{2k} = \frac{-v_0 \frac{R_2}{R_1 + R_2} U_1}{r_2}$$

$$R_{\text{aus}} = \frac{U_{21}}{I_{2k}} = \frac{-\frac{R_2}{R_1} U_1}{\frac{-v_0 \frac{R_2}{R_1 + R_2} U_1}{r_2}} = \frac{\frac{R_1 + R_2}{R_1} r_2}{v_0}$$

$$R_{\text{aus}} = \frac{\left(1 + \frac{R_2}{R_1}\right) r_2}{v_0} \approx \frac{|v|}{v_0} r_2$$

(3.22)

Nichtinvertierender Verstärker

Der nichtinvertierende Verstärker (Bild 3.31) weist einen sehr hohen Eingangswiderstand auf und wird deshalb auch als Elektrometervverstärker bezeichnet. Er bietet außerdem die Mög-

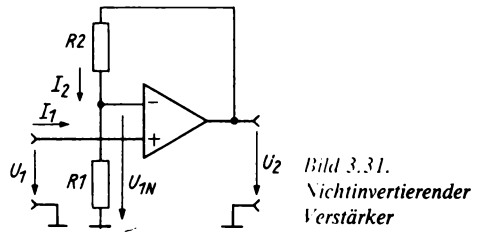


Bild 3.31.
Nichtinvertierender
Verstärker

lichkeit, die Verstärkung in weiten Grenzen zu variieren.

$$U_2 = U_1 \frac{R_1 + R_2}{R_1} = U_1 \left(1 + \frac{R_2}{R_1}\right)$$

$$v = \frac{U_2}{U_1} = 1 + \frac{R_2}{R_1}$$

(3.23)

Für den Eingangswiderstand ergibt sich

$$R_{\text{ein}} = v_0 \frac{r_1}{1 + \frac{R_2}{R_1}} \approx \frac{v_0}{v} r_1$$

(3.24)

(Beachte Hinweis im Abschn. 3.4.7.)

Der Eingangswiderstand des nichtinvertierenden Verstärkers vergrößert sich also gegenüber dem Eingangswiderstand des OPV um den Faktor

$$\frac{v_0}{1 + \frac{R_2}{R_1}}$$

und erreicht für $R_2 = 0$ ein Maximum. Der Ausgangswiderstand ist

$$R_{\text{aus}} = \frac{r_2}{v_0} \left(1 + \frac{R_2}{R_1} + \frac{R_2}{r_1}\right) \approx \frac{v}{v_0} r_2$$

(3.25)

Für $r_1 \rightarrow \infty$ folgt aus Gl. (3.25)

$$R_{\text{aus}} = \frac{r_2}{v_0} \cdot \frac{R_1 + R_2}{R_1} \quad (3.25a)$$

Für $v_0 \rightarrow \infty$ folgt aus Gl. (3.25)

$$R_{\text{aus}} = 0. \quad (3.25b)$$

Spannungsfolger

Der Spannungsfolger (Bild 3.32) hat einen sehr hohen Eingangs- und einen sehr niedrigen Ausgangswiderstand bei einer Spannungsverstärkung $v = 1$. Aus diesem Grund wird der Spannungsfolger u. a. als Impedanzwandler und zur Entkopplung von Netzwerken eingesetzt.

Der Spannungsfolger ist ein nichtinvertierender Verstärker mit $R_1 \rightarrow \infty$ und $R_2 = 0$. Nach Gl. (3.23) wird damit

$$U_2 = U_1 \text{ bzw. } v = \frac{U_2}{U_1} = 1 \quad (3.23 a)$$

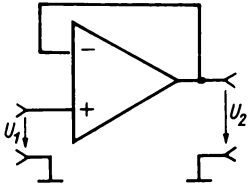


Bild 3.32.
Spannungsfolger

Der Eingangswiderstand ergibt sich nach Gl. (3.24)

$$R_{\text{ein}} = r_1 v_0 \quad (3.24 a)$$

(Beachte Hinweis im Abschn. 3.4.7)
und der Ausgangswiderstand nach Gl. (3.25)

$$R_{\text{aus}} = \frac{r_2}{v_0} \quad (3.25 c)$$

Die damit erreichbare Widerstandsübersetzung

$$\frac{R_{\text{ein}}}{R_{\text{aus}}} = \frac{r_1 v_0}{r_2 / v_0} = \frac{r_1}{r_2} v_0^2$$

ist also sehr hoch.

Differenzverstärker

Der Differenzverstärker (Bild 3.33) ist die Kombination eines invertierenden und eines nichtinvertierenden Verstärkers, mit dem die Differenz zweier auf Masse bezogenen Spannungen verstärkt wird.

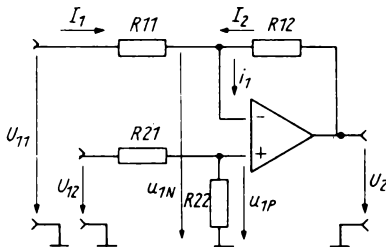


Bild 3.33. Differenzverstärker

Kenndaten eines Differenzverstärkers

Differenzverstärkung

$$v_D = \frac{U_2}{U_{12} - U_{11}} \quad (3.26)$$

Gleichtaktverstärkung

$$v_G = \frac{U_2}{U_G} \quad (3.27)$$

U_G Gleichaktspannung.

Werden die beiden Eingänge des Differenzverstärkers miteinander verbunden und mit einer gemeinsamen Spannung, der sog. Gleichaktspannung U_G , angesteuert, dann müßte die Ausgangsspannung $U_2 = 0$ sein. Das ist aus zwei Gründen nicht der Fall.

- Die notwendige Voraussetzung $R_{11}/R_{21} = R_{12}/R_{22}$ kann wegen der Widerstandstoleranzen praktisch nicht exakt erfüllt werden.
- Beim realen OPV wirkt an beiden Eingängen je ein Widerstand r_{1G} (Gleichtakteingangswiderstand) gegen Masse. Da $r_{1G} \gg r_1$, braucht dieser nur dann berücksichtigt zu werden, wenn die Gleichaktspannung sehr viel größer als die zwischen den Eingängen des OPV wirksame Differenzspannung ist. U_G darf einen von der Betriebsspannung abhängigen Höchstwert nicht überschreiten.

Gleichtaktunterdrückung

$$CMR = \frac{v_D}{v_G} \quad (3.28)$$

Die Gleichaktunterdrückung CMR ist ein Maß für die Güte des Differenzverstärkers. Im Idealfall ist $CMR \rightarrow \infty$; in der Praxis werden Werte von $10^3 \dots 10^5 = 60 \dots 100$ dB erreicht. CMR ist wesentlich von den Toleranzen der äußeren Schaltelemente abhängig. Die praktische Ausführung eines Differenzverstärkers muß sehr sorgfältig erfolgen, da sonst insbesondere bei kleinen Eingangsspannungen große Fehler auftreten können. Der auszuwählende OPV muß gute Gleichtakteigenschaften haben, und die Widerstände müssen enge Toleranzen aufweisen.

3.4.5.

Frequenzgang und seine Kompensation

Die hohe offene Spannungsverstärkung des OPV tritt nur bei sehr niedrigen Frequenzen auf. Mit höher werdenden Frequenzen sinkt wie bei jedem anderen Verstärker die Verstärkung v_0 ab, da die Kapazität der Ausgangstristoren des OPV an Einfluß gewinnen.

Der Verstärkungsabfall ist mit einer Phasendrehung verknüpft, die zur Instabilität der Schaltung führen kann.

Beim invertierenden Verstärker ist doch die Ausgangsspannung gegenphasig zur Eingangsspannung. Wird nun der auf den Eingang zurückgeführte Teil der Ausgangsspannung nochmals um $\varphi = -180^\circ$ gedreht und ist der Betrag der Schleifenverstärkung $|kv_0| \geq 1$, dann wird die Schwingbedingung erfüllt, und es kommt zur Selbsterregung. Der Zusammenhang zwischen Amplituden- und Phasengang soll zunächst an einem RC-Tiefpaß erläutert werden (Bild 3.34).

$$\left| \frac{U_2}{U_1} \right| = \frac{1}{\omega C} \cdot \frac{1}{\sqrt{R^2 + \frac{1}{\omega^2 C^2}}}$$

Für genügend hohe Frequenzen gilt

$$\left| \frac{U_2}{U_1} \right| \approx \frac{1}{\omega CR} = \frac{1}{f} \cdot \frac{1}{2\pi CR}$$

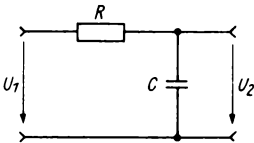


Bild 3.34. RC-Tiefpaß

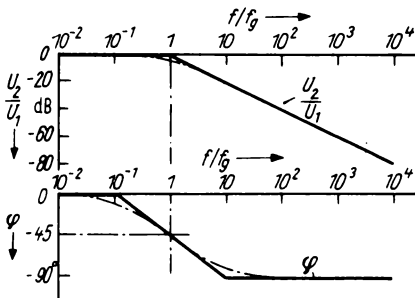


Bild 3.35. Amplituden- und Phasenfrequenzgang eines RC-Tiefpasses

Bezieht man f auf die Grenzfrequenz

$$f_g = \frac{1}{2\pi CR}, \text{ folgt}$$

$$\left| \frac{U_2}{U_1} \right| \approx \frac{1}{\frac{f}{f_g}}$$

Der zwischen U_2 und U_1 liegende Phasenwinkel ergibt sich aus

$$\tan \varphi = -\omega CR = -f \cdot 2\pi CR = -\frac{f}{f_g}$$

$$\text{Bei } \frac{f}{f_g} = 0 \quad \text{ist } \varphi = 0$$

$$\frac{f}{f_g} = 1 \quad \text{ist } \varphi = -45^\circ$$

$$\frac{f}{f_g} \rightarrow \infty \quad \text{ist } \varphi = -90^\circ$$

Die Darstellung des normierten Amplituden- und Phasenfrequenzgangs im doppelt logarithmischen Maßstab wird im Bild 3.35 gezeigt.

Dabei ist

$$\left| \frac{U_2}{U_1} \right|_{\text{dB}} = 20 \lg \left| \frac{U_2}{U_1} \right|$$

$$\text{Bei } \frac{f}{f_g} = 0 \quad \text{ist } \left| \frac{U_2}{U_1} \right| = 0 \text{ dB}$$

$$\frac{f}{f_g} = 1 \quad \text{ist } \left| \frac{U_2}{U_1} \right| = -3 \text{ dB}$$

$$\frac{f}{f_g} = 10 \quad \text{ist } \left| \frac{U_2}{U_1} \right| = -20 \text{ dB usw.}$$

Wächst f/f_g um das 10fache (1 Dekade = 1 dec), z. B. von 10 auf 100, dann vermindert sich $|U_2/U_1|$ auf 1/10 des bei $f/f_g = 10$ vorhandenen Wertes, also um

$$\frac{\left| \frac{U_2}{U_1} \right|_{100}}{\left| \frac{U_2}{U_1} \right|_{10}} = 20 \lg \frac{1}{10} = -20 \text{ dB.}$$

Bei genügend hohen Frequenzen fällt also $|U_2/U_1|$ um 20 dB/dec, während der Phasenwinkel $\varphi = -90^\circ$ beträgt.

Bei einem RC-Glied ist $\left| \frac{U_2}{U_1} \right| \sim \frac{1}{f}$, d. h. -20 dB/dec und $\varphi \approx -90^\circ$.

Bei zwei RC -Gliedern ist $\left| \frac{U_2}{U_1} \right| \sim \frac{1}{f^2}$, d. h. -40 dB/dec und $\varphi \approx -180^\circ$.

Bei drei RC -Gliedern ist $\left| \frac{U_2}{U_1} \right| \sim \frac{1}{f^3}$, d. h. -60 dB/dec und $\varphi \approx -270^\circ$.

Dieser Zusammenhang ist durch die Hilbert-Transformation gegeben, wonach allgemein für einen Abfall des Spannungsverhältnisses mit $1/f^n$ der Phasenwinkel $\varphi = -n\pi/2 = -n \cdot 90^\circ$ gilt.

Der prinzipielle Frequenzgang des Operationsverstärkers läßt sich damit leicht verstehen. Integrierte OPV sind meist 3stufig aufgebaut. Jede der drei Einzelstufen weist einen Frequenzgang auf, der vom Kollektorarbeitswiderstand und der Kollektorkapazität des Transistors bestimmt ist.

Der gesamte Amplituden- und Phasengang ergibt sich nun als Überlagerung der Amplituden- und Phasengänge der Einzelstufen mit ihren gegenseitig nicht verkoppelten RC -Gliedern. Damit keine ungünstigen Verhältnisse auftreten, müssen die Grenzfrequenzen der drei RC -Glieder sehr unterschiedlich sein. Die im doppelloarithmischen Maßstab dargestellte Verstärkung des OPV in Abhängigkeit von der Frequenz wird im Bild 3.36a und die Abhängigkeit des Phasenwinkels von der Frequenz im Bild 3.36b gezeigt. Diese Darstellungsform heißt *Bode-Diagramm*.

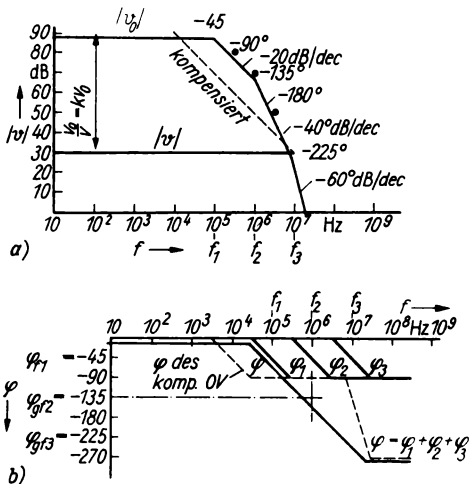


Bild 3.36. Bode-Diagramm des gegen- und nichtgegengekoppelten Operationsverstärkers

a) Amplitudenfrequenzgang; b) Phasenfrequenzgang

Mit steigender Frequenz sinkt die Verstärkung ab. Die Stufe mit der höchsten Kollektorzeitkonstante hat eine Grenzfrequenz f_1 , bei der $|v_0|$ um 3 dB abgefallen ist und der Phasenwinkel $\varphi_1 = -45^\circ$ beträgt. Der weitere Verstärkungsabfall erfolgt mit 20 dB/dec bis zur Frequenz f_2 , die die Grenzfrequenz der Stufe mit der mittleren Kollektorzeitkonstante ist. Bei f_2 hat der Phasenwinkel der zuvor betrachteten Stufe einen Wert von $\varphi_1 \approx -90^\circ$ erreicht, so daß sich nun eine Phase von $\varphi_{g2} \approx \varphi_1 + \varphi_2 = -90^\circ - 45^\circ = -135^\circ$ ergibt. Gleichzeitig geht ab f_2 der Verstärkungsabfall in 40 dB/dec über bis f_3 , der Grenzfrequenz der Stufe mit der niedrigsten Kollektorzeitkonstante. Bei f_3 ist der Phasenwinkel der beiden zuvor behandelten Stufen $\varphi_1 + \varphi_2 \approx -180^\circ$, so daß sich eine Phase $\varphi_{g3} = \varphi_1 + \varphi_3 = -180^\circ - 45^\circ = -225^\circ$ ergibt. Die Verstärkung fällt ab f_3 dann mit 60 dB/dec .

Wenn nun der Zusatzphasenwinkel $\varphi = -180^\circ$ beträgt und der Betrag der Schleifenverstärkung $|kv_0| = 1$ ist, tritt Selbsterregung ein. Um unter allen Umständen eine Selbsterregung zu vermeiden, bleibt man mit einer Phasensicherheit oder einem Phasenrand von 45° unter dem kritischen Wert von -180° , also bei $\varphi_{g2} = -180^\circ + 45^\circ = -135^\circ$. Die gewünschte Verstärkung v des gegengekoppelten Verstärkers verläuft bis zum Schnittpunkt mit der Frequenzkurve der offenen Spannungsverstärkung frequenzunabhängig, danach stimmt der Frequenzgang des gegengekoppelten Verstärkers mit demjenigen der offenen Spannungsverstärkung überein. Im Schnittpunkt ist $v = v_0$ bzw. $v_0/v = kv_0 = 1$.

Liegt nun der Schnittpunkt oberhalb der -135° -Frequenz, dann muß der Frequenzgang der offenen Spannungsverstärkung so korrigiert werden, daß bei der -135° -Frequenz der Verstärkungsabfall mit 20 dB/dec erfolgt. Je höher die Verstärkung v des gegengekoppelten Verstärkers, desto niedriger ist die Frequenz bei der $|v| = |v_0|$, desto schwächer kann also die Frequenzgangkorrektur sein.

Bild 3.36a zeigt die dargestellten Verhältnisse für $|v_0| = 90 \text{ dB}$, $|v| = 30 \text{ dB}$, Bild 3.36b für $f_{135^\circ} = 1 \text{ MHz}$. Die Frequenzgangkompensation ist in jedem Fall sehr sorgfältig vorzunehmen. Die Hersteller von OPV liefern dazu genaue Angaben, wo die Kompensationsglieder einzuschalten und wie diese je nach Rückkopplungsgrad bzw. Verstärkungsfaktor zu dimensionieren sind. Grundsätzlich unterscheidet man eine Eingangs- und eine Ausgangsfrequenzkompensation. Erstere ist für das Aussteuerverhalten

des Verstärkers günstig, bringt dafür aber einen Rauschanstieg. Vorteilhaft ist deshalb eine Verteilung der Kompensation auf Ein- und Ausgang.

Der Operationsverstärker A 109 ist wegen der hohen offenen Spannungsverstärkung und der intern wirkenden Gegenkopplung stets mit zwei Kompensationsgliedern zu beschalten, selbst wenn die äußere Gegenkopplung sehr klein oder Null wird.

Die Eingangskompensation erfolgt zwischen den Anschlüssen 3 und 12 durch ein $R_K C_{K1}$ -Glied, die Ausgangsfrequenzkompensation zwischen den Anschlüssen 9 und 10 durch einen Kondensator C_{K2} .

Beim Operationsverstärker A 109 wird vom Hersteller empfohlen, daß in den Ausgang ein Widerstand $R_4 = 51 \Omega$ eingeschaltet wird, um eine unerwünschte Schwingung zu vermeiden (Bild 3.37). Die Diode V dient als Schutz gegen eine mögliche Blockierung (Latch up; Abschn. 3.4.7.)

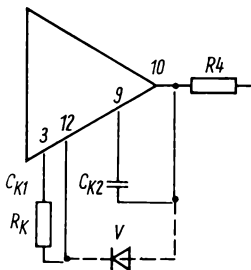


Bild 3.37. Frequenzkompensation des A 109

Die Frequenzkompensation bei den OPV-Schaltkreisen der 2. Generation ist wesentlich einfacher zu realisieren. Sie erfolgt bei einigen

Tafel 3.5. Kompensationsglieder des A 109

v dB	C_{K1} pF	C_{K2} pF	R_K k Ω
0	4700	200	1,5
5	2700	100	1,5
10	1500	62	1,5
14	1000	39	1,5
20	470	20	1,5
25	270	10	1,5
30	150	6	1,5
40	100	3	1,5
50	27	3	1,5
≥ 60	10	3	0

Anmerkung: Bei hohen Verstärkungen müssen u. U. etwas höhere Werte für C_{K1} gewählt werden, um die Phasendrehung des Gegenkopplungszweiges auszugleichen.

Typen intern, bei anderen durch einen einzigen extern anzuschließenden Kondensator (Bilder 3.43 bis 3.48).

3.4.6.

Nullpunktfehler und seine Kompensation

Die Ausgangs-Eingangs-Kennlinie des OPV verläuft im Idealfall durch den Nullpunkt mit einem linearen Anstieg von v (Bild 3.38).

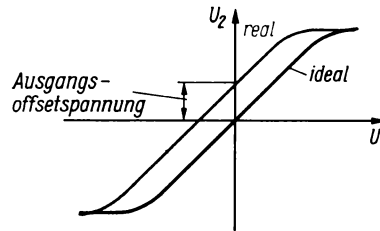


Bild 3.38. Ausgangs-Eingangs-Kennlinie des Operationsverstärkers

Real stellt sich aber immer eine positive oder negative Ausgangsspannung ein, obwohl die Eingangsspannung Null ist. Diese von Null abweichende Ausgangsspannung heißt Ausgangs-offsetspannung. Sie wird durch folgende Größen verursacht:

- Eingangsruhestrome I_N und I_P (Biasstrom)
- Unsymmetrien innerhalb des OPV
- Drift (Änderung) der genannten Größen, die durch Temperatur- und Betriebsspannungsschwankungen und Alterung der Bauelemente auftreten.

Alle Maßnahmen, die dazu dienen, den Nullpunktfehler zu beseitigen, werden unter dem Begriff *Offsetkompensation* zusammengefaßt.

Zur Kompensation der Wirkung der Eingangsruhestrome eignen sich die Schaltungen nach Bild 3.39. Da I_N und I_P an den Widerständen

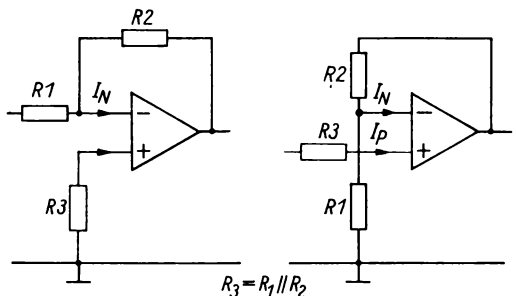


Bild 3.39. Kompensation der Wirkung der Eingangsruhestrome

des Eingangskreises Spannungsabfälle verursachen, wird durch deren Differenz der OPV in unerwünschter Weise angesteuert. Die Wirkung läßt sich durch einen Widerstand R_3 kompensieren, wenn $R_3 = R_1 \parallel R_2$ ist. Der durch Unsymmetrien des OPV entstehende Nullpunktfehler kann als innere Störspannungs-

quelle im OPV aufgefaßt werden. Legt man eine kleine, von der Betriebsspannung abgeleitete Gegenspannung an den Eingang, dann kann der Offset vollständig kompensiert werden. Am häufigsten werden die im Bild 3.40 angegebenen Kompensationsschaltungen angewandt.

Eine weitere Möglichkeit der Offsetkompensation ist bei bestimmten OPV-Typen durch ein extern anzuschließendes Potentiometer möglich, mit dem der Symmetriefehler des Stromspiegels der Eingangsstufe korrigiert werden kann (z. B. bei den Typen B 061, B 066, B 081, B 083, B 176, B 177).

Der durch Drift verursachte Einfluß wird durch die angegebenen Schaltungen ebenfalls vermindert, läßt sich aber durch einfachen Abgleich nicht restlos beseitigen. Dafür sind, wenn erforderlich, spezielle Maßnahmen erforderlich.

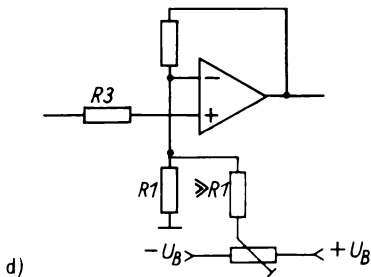
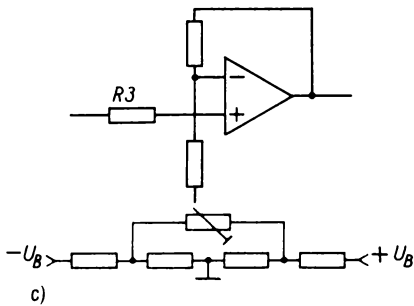
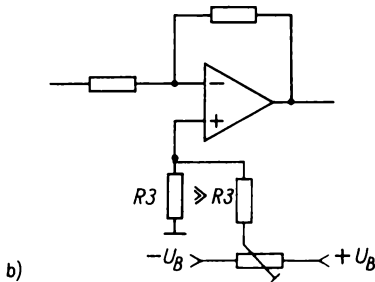
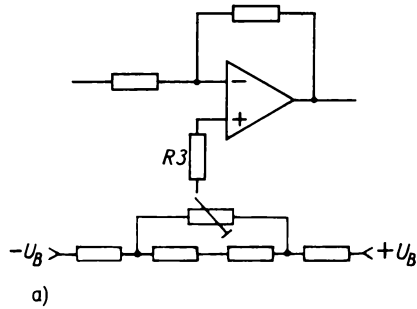


Bild 3.40. Schaltungen zur Offsetkompensation

a) und b) beim invertierenden Operationsverstärker

c) und d) beim nichtinvertierenden Operationsverstärker

3.4.7.

Kenngrößen von Operationsverstärkern

Die Kenngrößen von OPV werden unterteilt in

- Grenzwerte
- Betriebswerte
- Statische Kennwerte
- Dynamische Kennwerte.

Grenzwerte sind Maximalwerte, die keinesfalls (auch nicht kurzzeitig) überschritten werden dürfen, weil sonst der Schaltkreis zerstört wird oder seine Eigenschaften bleibend verschlechtert werden. Für den Anwender sind vor allem die Grenzwerte folgender Größen wichtig:

- Betriebsspannung
- Verlustleistung
- Umgebungstemperatur
- Ausgangsgleichstrom (Zerstörgrenze für die Ausgangsstufe)
- Gleitakteingangsspannung (Zerstörgrenze für die Eingangsstufen)
- Differenzeingangsspannung (Zerstörgrenze, hinter der der Durchbruch des in Sperrichtung betriebenen Eingangstransistors erfolgt).

Betriebswerte sind die Werte, bei denen der Schaltkreis betrieben werden darf und funktionsfähig ist. Oft werden Wertebereiche angegeben. Es gehören dazu

- Betriebsspannungsbereich
- Betriebstemperaturbereich
- Stromaufnahme (Ruhestromaufnahme) bzw. Ruheverlustleistung
- Eingangsspannungsbereich

- Aussteuerbereich der Ausgangsspannung; die über den Endstufentransistoren entstehenden Spannungsabfälle begrenzen die minimale und maximale Ausgangsspannung. Diese Spannungen hängen bei ausreichend hoher Übersteuerung nur von der Betriebsspannung und dem Belastungswiderstand ab und werden unter diesen Randbedingungen angegeben.

Statische Kennwerte. Diese gelten für bestimmte festgelegte bzw. gewählte Betriebswerte. Da sie alle auch von der Umgebungstemperatur abhängig sind, werden häufig die Temperaturkoeffizienten der Offsetspannung und des Offsetstromes angegeben.

Spannungsverstärkung in offener Schleife (offene Spannungsverstärkung; Leerlaufspannungsverstärkung) v_0 . Dieser Kennwert wird oft als der wichtigste angesehen. Er gilt für den nicht rückgekoppelten OPV für tiefe Frequenzen:

$$v_0 = \frac{u_2}{u_D} \text{ mit } u_D = u_{IP} - u_{IN}$$

Größenordnung: $v_0 > 10^4 = 80 \text{ dB}$

Eingangswiderstand r_1 . Der Eingangswiderstand setzt sich zusammen aus dem Differenzeingangswiderstand r_{1D} (Widerstand zwischen den beiden Eingangsklemmen) und dem Gleichakteingangswiderstand r_{1G} (Widerstand zwischen beiden miteinander verbundenen Eingangsklemmen und Masse). Letzterer liegt in der Größenordnung $r_{1G} > 30 \text{ M}\Omega$, ist also viel größer als r_{1D} . Da beim nichtinvertierenden Betrieb infolge der Seriengegenkopplung der Differenzeingangswiderstand um den Faktor v_0/v vergrößert wird, ist der parallel liegende Gleichakteingangswiderstand nicht mehr zu vernachlässigen und begrenzt den für die Schaltung möglichen Eingangswiderstand $R_{\text{ein}} = \frac{v_0}{v} r_{1D} \parallel r_{1G}$.

Ausgangswiderstand r_2 . Das ist der Widerstand zwischen der Ausgangsklemme und Masse, wenn beide Eingänge auf Masse liegen.

Eingangsbiasstrom I_B . Zur Arbeitspunkteinstellung des realen OPV fließen die statischen Eingangsströme I_N und I_P . Bei bipolaren Eingangsstufen sind das die Basisströme der Eingangstransistoren und haben eine Größenordnung von nA bis μA . Bei Eingangsstufen mit Sperrschichtfeldeffekttransistoren haben sie eine Größe von pA bis nA. Der Mittelwert beider

Eingangsruheströme wird als Biasstrom I_B bezeichnet.

$$I_B = \frac{I_N + I_P}{2}$$

Die Ruheströme verursachen an den Widerständen des Eingangskreises Spannungsabfälle, deren Differenz eine unerwünschte Aussteuerung des OPV zur Folge hat. Diese Spannungsabfälle müssen deshalb mit einer Schaltung nach Bild 3.39 kompensiert werden.

Eingangsoffsetstrom I_{off} . Da sich beim realen OPV die Eingangsruheströme I_N und I_P nur wenig unterscheiden, ist es zweckmäßig, ihre Differenz zu erfassen. Diese wird als Offsetstrom I_{off} bezeichnet:

$$I_{\text{off}} = |I_N - I_P|$$

Eingangsoffsetspannung U_{off} . Durch unvermeidliche Unsymmetrien innerhalb des OPV erhält man eine Ausgangsspannung, obwohl die Eingangsdifferenzspannung Null ist. Man bezeichnet nun die Spannung, die zwischen die Eingangsklemmen gelegt werden muß, damit die Ausgangsspannung Null wird, als Eingangsoffsetspannung (Eingangsfehlerspannung) U_{off} . Die erforderliche Kompensation kann nach den Schaltungen entsprechend Bild 3.40 erfolgen.

TK der Eingangsoffsetspannung
 $\Delta U_{\text{off}}/\Delta \theta$ (ca. $5 \mu\text{V/K}$)

TK des Eingangsoffsetstromes
 $\Delta I_{\text{off}}/\Delta \theta$ (ca. $0,5 \text{ nA/K}$)

Gleichtaktunterdrückung CMR (common-mode-rejection). Erwünscht ist, daß die Ausgangsspannung nur von der Differenzeingangsspannung beeinflusst wird. Eine Gleichtakteingangsspannungsänderung sollte also die Ausgangsspannung nicht verändern. Das ist real nicht erreichbar, da die Gleichtaktverstärkung zwar sehr klein, aber nicht Null ist.

$$CMR = \frac{v}{v_G}$$

Eine hohe Gleichtaktunterdrückung des OPV ist immer erwünscht. Typisch sind für $CMR > 60 \dots 100 \text{ dB}$.

Bei invertierenden Verstärkern tritt praktisch keine Änderung der Gleichtakteingangsspannung auf, die endliche CMR ist also vor allem bei nichtinvertierenden Verstärkern zu beachten.

Betriebsspannungsunterdrückung SVR (supply-voltage-rejection). Bei diesem Kennwert handelt es sich um die Unterdrückung der Betriebsspannungsänderung. Da letztere auch eine Änderung der Ausgangsspannung verursacht, muß dieser Einfluß beseitigt werden, z. B. durch stabilisierte Betriebsspannungen und für Netzstörungen (vor allem höherfrequenter Anteile) durch RC -Siebglieder in den Versorgungsleitungen. SVR ist definiert als Verhältnis der notwendigen Differenzeingangsspannungsänderung zur Betriebsspannungsänderung, damit die Ausgangsspannung konstant bleibt. Typisch sind für

$$SVR < 200 \mu V/V \text{ bzw. } SVR > 74 \text{ dB.}$$

Dynamische Kennwerte. Diese gelten ebenfalls für bestimmte festgelegte bzw. gewählte Betriebswerte.

Anstiegszeit t_r . Bei Übersteuerung einer oder mehrerer Stufen des OPV z. B. mit einem Spannungssprung vergeht wegen der internen Kapazitäten eine bestimmte Zeit, bis das Ausgangssignal den Endwert erreicht. Die Anstiegszeit t_r (genauer Flankenanstiegszeit) ist die Zeit, in der das Ausgangssignal von 10 % auf 90 % des Endwertes angestiegen ist.

Slew-rate SR ist die maximale Änderungsgeschwindigkeit bzw. Flankensteilheit der Ausgangsspannung des OPV bei Übersteuerung mit großen Rechteckimpulsen am Eingang. Die Slew-rate ergibt sich durch die endliche Umladezeit der internen Kapazitäten und der externen Kompensationskondensatoren. SR liegt in der Größenordnung von 0,2 bis 100 V/ μ s.

Die Slew-rate ist ein aussagekräftiger Kennwert bezüglich der Aussteuerbarkeit und des nutzbaren Frequenzbereiches. Erfolgt beispielsweise die Aussteuerung mit einer sinusförmigen Größe, dann muß zur Aufrechterhaltung einer bestimmten Ausgangssignalamplitude mit zunehmender Frequenz die Eingangssignalspannung erhöht werden, da die offene Spannungsverstärkung des OPV bei höheren Frequenzen abfällt. Das führt schließlich zur Übersteuerung und zu einer durch SR begrenzten Anstiegsgeschwindigkeit der Ausgangsspannung. Das Ausgangssignal bleibt unverzerrt, solange die Bedingung $\hat{U}_2 \cdot 2\pi f \leq SR$ eingehalten wird.

Beispiel

Am Ausgang eines OPV soll eine unverzerrte sinusförmige Spannung mit einer Amplitude $\hat{U}_2 = 5 \text{ V}$ und einer Frequenz $f = 50 \text{ kHz}$ bereitgestellt werden.

Die Anstiegsgeschwindigkeit dieser Spannung ist dann

$$\hat{U}_2 \cdot 2\pi f = 5 \text{ V} \cdot 2 \cdot 3,14 \cdot 50 \text{ kHz} \approx 1,571 \text{ V}/\mu\text{s.}$$

Es ist also mindestens eine Slew-rate

$$SR \geq 1,571 \text{ V}/\mu\text{s} \text{ erforderlich.}$$

Es gibt weitere statische und dynamische Kennwerte, die bei Bedarf den Bauelementenunterlagen zu entnehmen sind. Zu erwähnen ist noch der Latch-up-Effekt. Er ist kein Kennwert, führt nicht zur Zerstörung des Schaltkreises, aber zur Funktionsunfähigkeit. Latch-up bedeutet Blockierung. Wird nämlich der OPV-Ausgang von der Eingangsspannung in den Sättigungsbereich getrieben, dann kann er in diesem Zustand auch bei Wegfall dieser Eingangsspannung „hängenbleiben“. Mehrere OPV sind mit einem internen Latch-up-Schutz ausgerüstet (Latch-up-frei). Ist das nicht der Fall, sind eventuell externe Schaltungsmaßnahmen notwendig. Für den Typ A 109 kann das durch eine Diode zwischen pin 10 und pin 12 entsprechend Bild 3.37 erfolgen.

3.4.8.

Eigenschaften und Schaltungstechnik verschiedener Operationsverstärkertypen

Integrierte Operationsverstärker haben eine eminente Bedeutung erlangt. Sie sind universell einsetzbar, und es lassen sich damit zahlreiche lineare und nichtlineare Schaltungsprobleme lösen. Entsprechend der verschiedenen Aufgaben wurde international ein breites Typenspektrum der OPV entwickelt.

Während die universellen, für allgemeine Anwendung geeigneten OPV der 1. Generation vorwiegend npn-Strukturen und zahlreiche Widerstände enthalten (Typ A 109), wird bei den OPV-Schaltkreisen der 2. Generation ein komplementäres Schaltungskonzept mit npn- und pnp-Strukturen und Sperrschichtfeldeffekttransistoren mit wenig integrierten Widerständen angewandt.

Bild 3.41 zeigt die Entwicklung und Einteilung der integrierten OPV. Die OPV der 2. Generation weisen gegenüber der 1. oft verbesserte Eigenschaften auf, wie beispielsweise höhere Grenzwerte für die Differenz- und Gleichtakteingangsspannung, kein Blockieren (Latch-up-frei), einfache Offsetkompensation, einfache externe, u. U. interne Frequenzkompensation, geringe Leistungsaufnahme, kurzschlußfeste Endstufe, TTL-kompatibel. Bei programmierbaren Typen lassen sich verschiedene Parameter des

OPV der
1. Generation

OPV der 2. Generation

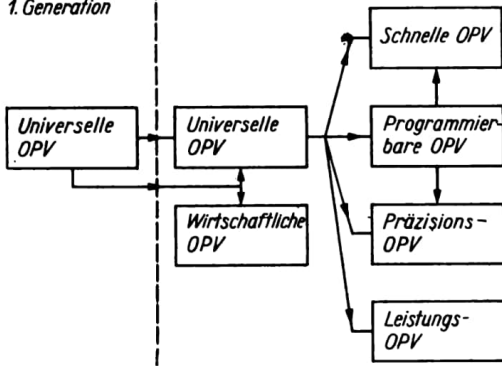


Bild 3.41. Einteilung der Operationsverstärker

Schaltkreises von außen steuern; Leistungs-OPV haben eine Gesamtverlustleistung von über 20 W und sind für Ausgangsspitzenströme von 3,5 A ausgelegt.

Im folgenden werden einige OPV-Schaltkreisfamilien vorgestellt, die der VEB Halbleiterwerk Frankfurt (Oder) fertigt. Neben den wichtigsten Eigenschaften und Anschlußbedingungen werden auch einige Innenschaltungen bzw. deren Übersichtspläne wiedergegeben, aus denen der Anwender wichtige Erkenntnisse für spezielle Einsatzfälle gewinnen kann.

Der Schaltkreis A 109. Dieser universelle OPV-Schaltkreis der 1. Generation besteht aus drei Stufen (Bild 3.42).

Die erste Stufe (Eingangsstufe) ist ein mit den Transistoren $V1$ und $V2$ gebildeter Differenzverstärker mit der Konstantstromquelle $V10$, $V11$. Letztere spiegelt den Strom auf etwa $40 \mu A$. Auf Grund der niedrigen Kollektorströme der Eingangstransistoren kommt es zu einem relativ hohen Eingangswiderstand von $> 50 k\Omega$ des OPV. Die zweite Stufe ist ein in Darlington-Schaltung ausgeführter Differenz-

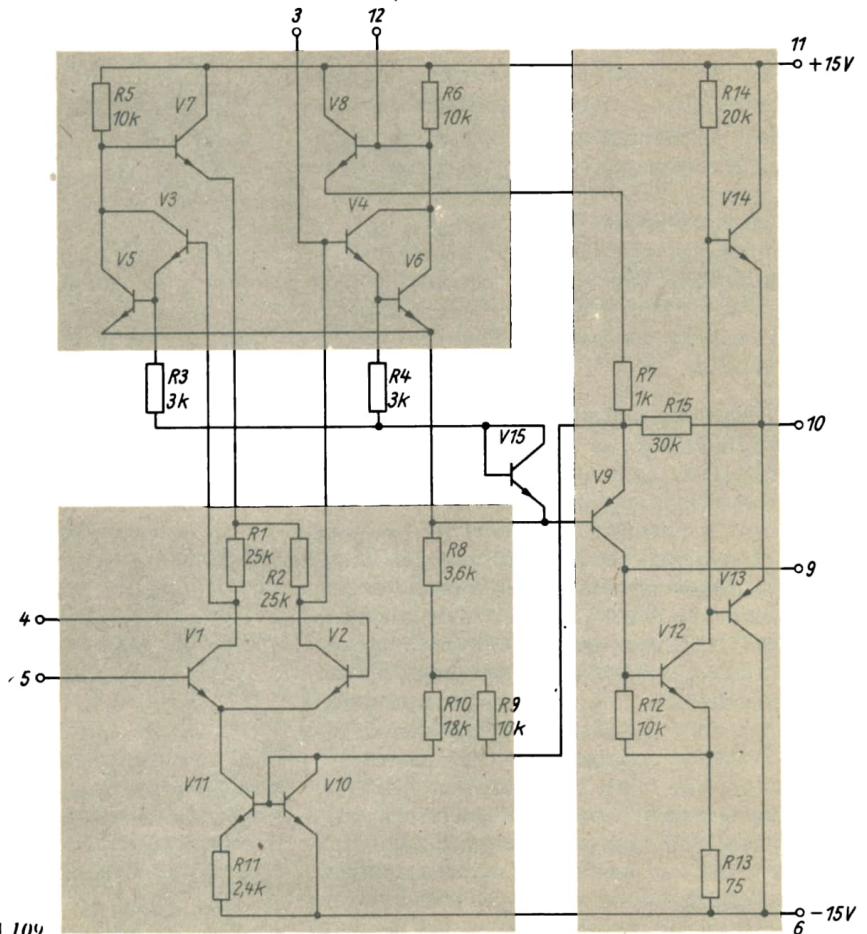


Bild 3.42. Innenschaltung des Operationsverstärkers A 109

verstärker mit den Transistoren $V3$, $V5$ und $V4$, $V6$. Der Transistor $V7$ entlastet den Widerstand $R5$ vom Summenstrom der ersten Stufe, der in Kollektorschaltung arbeitende Transistor $V8$ führt die Ausgangsspannung der zweiten Stufe an den Eingang der Endstufe. Durch eine Temperaturkompensation mit $R3$, $R4$ und $V15$ wird eine sehr geringe Temperaturabhängigkeit erreicht. Die zweite Stufe weist eine große Spannungsverstärkung auf und ermöglicht den Übergang vom symmetrischen Eingang zum unsymmetrischen Ausgang.

Die dritte Stufe ist die Ausgangsstufe. Der pnp-Transistor $V9$ dient der Pegelverschiebung, um einen großen Aussteuerbereich der Ausgangsspannung zu erreichen. Die Treiberstufe mit $V12$ erbringt die erforderliche Spannungsverstärkung, um die komplementäre Gegentakt-B-Endstufe mit dem pnp-Transistor $V13$ und dem npn-Transistor $V14$ auszusteuern. Im Interesse geringer Verzerrungen und eines kleinen Ausgangswiderstands ist der Gegenkopplungswiderstand $R15$ eingefügt.

Neben dem Vorteil des universellen Einsatzes hat der A 109 u. a. folgende Nachteile

- nicht kurzschlußfest; die Dauer des Kurzschlußausgangsstromes darf maximal 5 s betragen
- nicht latch-up-frei
- aufwendige externe Frequenzkompensation durch eine Eingangs- und Ausgangskompensation; es sind also stets zwei Kompensationsglieder entsprechend dem Bild 3.37 erforderlich.

Die Schaltkreisfamilien B 761 und B 861. Diese OPV-Schaltkreise der 2. Generation sind universell einsetzbar und besonders kostengünstig. Die Schaltung zeigt Bild 3.43. Die Eingangsstufe ist als einfacher Differenzverstärker ausgeführt. Das verstärkte Signal wird an den beiden Kollektorwiderständen abgegriffen und gelangt an eine zweite Verstärkerstufe, in der die Funktionen Invertieren und Nichtinvertieren realisiert werden. Der Endverstärker arbeitet als Darlingtonschaltung. Eine Konstantspannungsquelle erzeugt eine stabile Spannung für die Konstantstromversorgung der Verstärkerstufen. Dadurch wird eine niedrige untere Betriebsspannungsgrenze, eine gute Betriebsspannungsunterdrückung und ein geringer Betriebsspannungseinfluß auf andere Kenngrößen erreicht. Der Ausgang ist als offener Kollektorausgang (open-collector) realisiert, über dem der Aus-

gangsstrom bis 70 mA eingestellt werden kann. Der open-collector-Ausgang ermöglicht im Vergleich zu Gegentaktendstufen ein lineares Übertragungsverhalten.

Die Gleichakteingangsspannung darf maximal gleich der Betriebsspannung sein; ebenso die Differenzeingangsspannung, letztere aber höchstens ± 15 V. Der Betriebsspannungsbereich liegt bei B 761 zwischen $\pm 1,5$ bis ± 18 V, bei B 861 zwischen $\pm 1,5$ bis ± 10 V. Die Ergänzungstypen B 765 und B 865 sind bei gleichen Kennwerten für den erweiterten Temperaturbereich von -25 bis $+85^\circ\text{C}$ geeignet.

Wegen des open-collector-Ausgangs muß in jedem Fall ein Lastwiderstand R_a gegen die positive Betriebsspannung geschaltet werden. Dieser

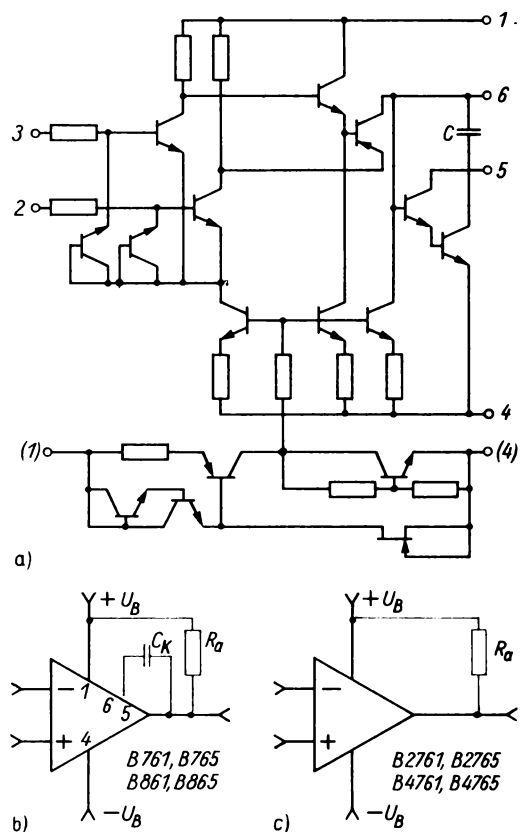


Bild 3.43. OPV mit einfachem Differenzverstärker-Eingang, Darlington-Ausgang und externer (bzw. interner) Frequenzgangkompensation

- a) Innenschaltung; die integrierte Kapazität $C \approx 18$ pF ist nur bei den Doppel- und Vierfach-OPV enthalten
- b) Anschlußschema der einfachen OPV
- c) Anschlußschema der Doppel- und Vierfach-OPV mit interner Frequenzkompensation

ist so zu dimensionieren, daß der zulässige Ausgangsstrom von 70 mA nicht überschritten wird. Für einfache Anwendungen mit nicht zu großem Ausgangsstrom und einer Betriebsspannung von ± 15 V kann $R_a = 2$ k Ω betragen. Die externe Frequenzkompensation erfolgt durch einen Kondensator C_K entspr. der Schaltung nach Bild 3.43b. Je kleiner die Kapazität dieses Kondensators gewählt wird, desto größer wird die Bandbreite. C_K sollte aber stets ≥ 3 pF sein, da sonst die Schaltung schwingen kann. Für die Offsetkompensation gibt es keine interne Abgleichmöglichkeit. Muß eine Offsetkompensation erfolgen, kann diese durch eine im Bild 3.40 angegebene Schaltung erfolgen.

Die Schaltkreisfamilien B 2761 und B 4761. Der OPV B 761 ist im Schaltkreis B 2761 zweimal in einem 8poligen DIL-Gehäuse (Doppel-OPV), im Schaltkreis B 4761 viermal in einem 14poligen DIL-Gehäuse (Vierfach-OPV) enthalten. Sowohl beim Doppel- wie beim Vierfach-OPV erfolgt die Frequenzkompensation intern durch eine Kapazität $C \approx 18$ pF in jedem Kanal. Ein äußerer Anschluß fehlt, so daß eine externe Kompensation entfällt. Das Einsatzgebiet dieser Schaltkreise ist damit eingeschränkt; bevorzugte Anwendungen liegen in der NF-Technik und als Gleichspannungsverstärker. Die Ergänzungstypen B 2765 und B 4765 haben einen erweiterten Temperaturbereich von -25 bis $+85$ °C.

Der Schaltkreis B 631. Dieser ist wie die OPV B 761 und B 861 aufgebaut, nur daß anstelle der einfachen Differenzverstärker-Eingänge ein Darlington-Differenzverstärker-Eingang vorliegt. Damit wird ein größerer Eingangswiderstand von $r_i > 3$ M Ω und ein entsprechend kleiner Eingangsbiassstrom $I_B < 50$ nA erreicht. Die Eingangsdifferenzspannung darf aber ± 13 V nicht überschreiten. Der Ergänzungstyp B 635 ist für den erweiterten Temperaturbereich einsetzbar. Innenschaltung und Anschlußschema des B 631 und B 635 zeigt Bild 3.44.

Die Schaltkreise B 621 und B 611. Beide Schaltkreise (einschließlich die für erweiterten Temperaturbereich gefertigten Ergänzungstypen B 625 und B 615) haben einen TTL-kompatiblen Ausgang. Dabei ist die Darlingtonendstufe aufgetrennt und beide Kollektoren sind nach außen geführt. Die Restspannung der OPV ist dadurch sehr klein und entspricht einer U_{CE} -

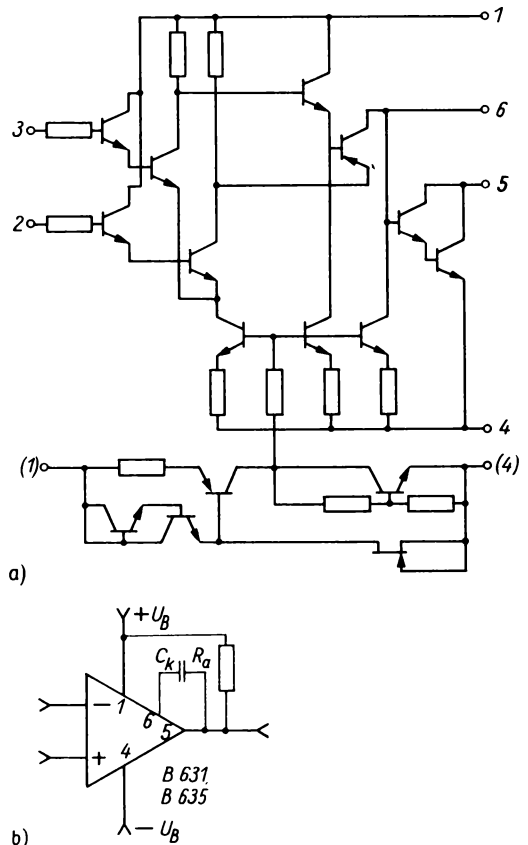


Bild 3.44. OPV mit Darlington-Differenzverstärker-Eingang, Darlington-Ausgang und externer Frequenzkompensation

a) Innenschaltung; b) Anschlußschema

Sättigungsspannung. Der Low-Pegel für TTL-Schaltkreise wird damit auch bei niedrigen Betriebsspannungen erreicht. Eine direkte Ansteuerung der TTL-Schaltkreise ist somit möglich.

Zu beachten ist, daß von dem zweiten nach außen geführten Kollektor ebenfalls ein Widerstand zur positiven Betriebsspannung geschaltet werden muß. Dieser sollte $R = (3 \dots 20)R_a$ betragen. Der Strom darf aber 10 mA nicht überschreiten.

Für die Bemessung von R_a gelten die gleichen Bedingungen wie für die anderen OPV-Typen mit Darlingtonendstufe.

Die für eine bestimmte Schaltung gewählten Werte R_a und R richten sich nach dem geforderten Ausgangsstrom. Denn dieser beeinflusst auch die Größe und den Verlauf der Restspannung, d. h. den Aussteuerbereich.

Der Schaltkreis B 621 (B 625) hat einen einfachen Differenzverstärker-Eingang, der B 611 (B 615) einen Darlington-Differenzverstärker-Eingang. Innenschaltungen und Anschlußschemata zeigen die Bilder 3.45 und 3.46.

Die OPV mit TTL-kompatiblen Ausgang haben keine Anschlußmöglichkeit für eine Frequenzkompensation. Sie schwingen deshalb bei kleinen Verstärkungen und arbeiten erst bei $v \geq 60$ dB stabil. Sie werden deswegen bevorzugt als Schwellwertschalter (Schmitt-Trigger) und Komparatoren eingesetzt.

Die Schaltkreisfamilien B 080 und B 060 (BIFET-Operationsverstärker)

Diese Schaltkreise haben einen bipolaren Aufbau, ihre Eingangsdifferenzverstärker sind aber mit Sperrschichtfeldeffekttransistoren (SFET) realisiert worden. Dadurch wird ein sehr hoher Eingangswiderstand $r_1 \approx 10^{12} \Omega$ bei sehr kleinen

Bias- und Offsetströmen erreicht. Weitere Eigenschaften sind geringe Leistungsaufnahme, sie sind latch-up-frei, haben große Bereiche für die Differenz- und Gleichtaktingangsspannungen und sind kurzschlußfest bei Einhaltung der maximalen Verlustleistung. Durch die interne Frequenzkompensation ist nur eine geringe Außenbeschaltung notwendig.

Die B 060er Typen unterscheiden sich von den B 080er durch noch geringere Leistungsaufnahme (low power), eignen sich deshalb besonders für Geräte mit Batteriebetrieb. Diesem Vorteil stehen als Nachteil die etwas schlechteren dynamischen Kennwerte (kleinere Slewrate, kleinere Bandbreite, größeres Rauschen) gegenüber. Beide Schaltkreisfamilien haben eine zwar im Detail abweichende, im Grundsätzlichen aber ähnliche Schaltungskonfiguration und sind (außer dem B 066) pinkompatibel.

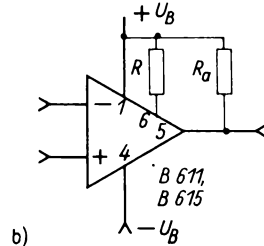
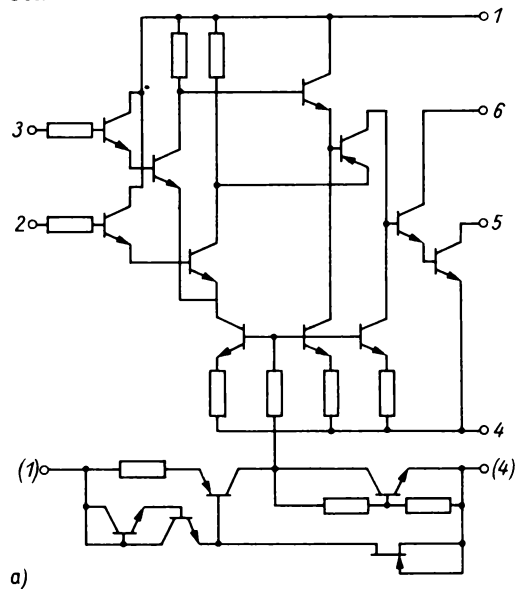
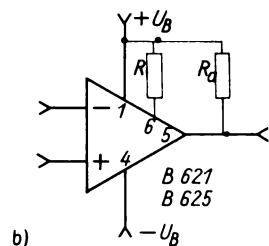
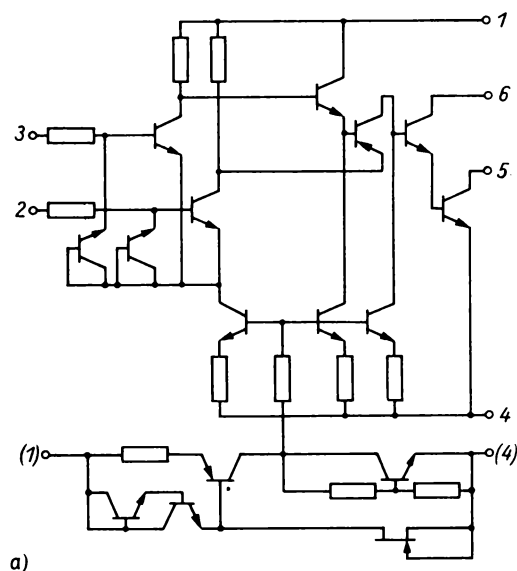


Bild 3.45. OPV mit einfachem Differenzverstärker-Eingang, TTL-kompatiblen Ausgang, ohne Frequenzkompensation

a) Innenschaltung; b) Anschlußschema

Bild 3.46. OPV mit Darlington-Differenzverstärker-Eingang, TTL-kompatiblen Ausgang, ohne Frequenzkompensation;

a) Innenschaltung; b) Anschlußschema

Innerhalb der Schaltkreisfamilien gibt es folgende OPV-Typen:

- Einfach-OPV (B 081, B 061)
- Doppel-OPV (B 082, B 062)
- Vierfach-OPV (B 084, B 064)

Ferner die Sondertypen

- Einfach-OPV mit externer Frequenzkompensation (B 080, B 060)
- Doppel-OPV mit Offsetkompensationsanschlüssen (B 083)
- Einfach-OPV mit der Möglichkeit einer externen Leistungsaufnahmesteuerung (B 066).

Bild 3.47 zeigt die Innenschaltung der B 080er Schaltkreise (a), und die Schaltungen zur Frequenz- und Offsetkompensation der B 080er und B 060er Typen (b-d).

Der Präzisionsoperationsverstärker B 087

Aufbau und Kennwerte dieses Schaltkreises sind ähnlich dem des B 081. Der B 087 ist aber, besonders rauscharm und weist eine sehr kleine Drift der Offsetgrößen auf. Präzisions-OPV werden dann eingesetzt, wenn spezielle, meistens sehr hohe Anforderungen an eine Schaltung gestellt werden. Das ist vor allem bei kommerziellen Geräten der Fall.

Die Schaltkreise B 176 und B 177 sind programmierbare Kleinleistungs-Operationsverstärker und grundsätzlich für alle Schaltungen der typischen Operationsverstärkertechnik geeignet. Der Vorteil programmierbarer OPV gegenüber nicht programmierbaren besteht darin, daß bei ihnen viele Parameter geändert werden können. Das geschieht durch einen als Setstrom bezeichneten Steuerstrom, der zwischen 1,5 und 200 μA liegen kann. Der eingespeiste Setstrom ist der Basisstrom einer internen pnp-Strombank, mit dem die Arbeitspunkte der einzelnen OPV-Stufen eingestellt werden. Mit steigendem Setstrom werden einige Parameter günstiger, andere ungünstiger. Der Anwender muß also einen Kompromiß eingehen und einen solchen Arbeitspunkt wählen, bei dem die Mehrzahl der Parameter die für den konkreten Einsatzfall zweckmäßigsten Werte ergeben. Prinzipiell steigen r_{in} größer werdendem Setstrom auch der Eingangsbiasstrom, die Ruhestromaufnahme, die offene Spannungsverstärkung und die Slewrate. Für viele Fälle werden optimale Werte bei einem mittleren Setstrom von 15 μA erreicht.

Die Schaltkreise B 176 und B 177 können mit einer Betriebsspannung von $\pm 1,2$ bis $\pm 15\text{ V}$ betrieben werden. Der Typ B 176 ist im Gegen-

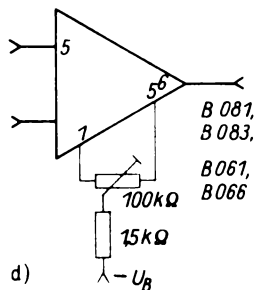
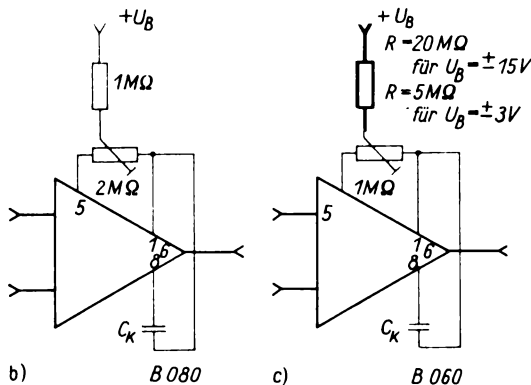
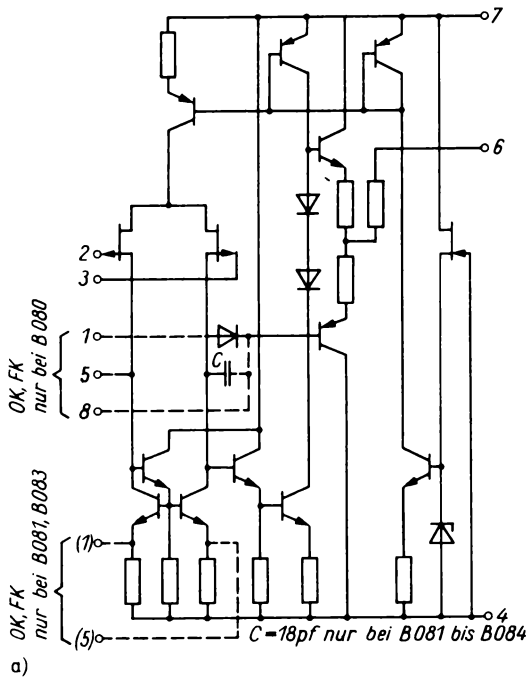


Bild 3.47. BIFET-Operationsverstärker

a) Innenschaltung der Schaltkreise B 080 bis B 084
b) bis d) Frequenzgang- und Offsetkompensation der Schaltkreisfamilien B 080 und B 060

satz zum B 177 mit einer Kapazität $C = 30 \text{ pF}$ intern frequenzkompensiert.

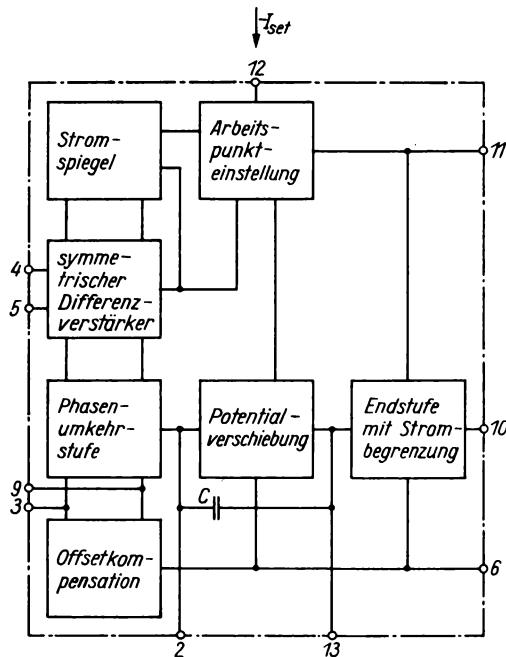
Bild 3.48 zeigt den Übersichtsplan der Innenschaltung der programmierbaren OPV (a) und Schaltungsmöglichkeiten für die Einspeisung des Setstromes (b und c).

Der Schaltkreis B 165 ist ein Leistungsoperationsverstärker. Leistungsoperationsverstärker werden benötigt als Schaltverstärker für Lampen, Motoren u. ä., als Leistungs-Interface, als Strom- und Spannungsquellen und als Impuls- und Signalgeneratoren großer Leistung.

Der B 165 kann mit einer Betriebsspannung von ± 6 bis $\pm 18 \text{ V}$ betrieben werden, läßt einen Ausgangsspitzenstrom von $3,5 \text{ A}$ und einen Ausgangsleichstrom von $2,5 \text{ A}$ zu bei einer Ge-

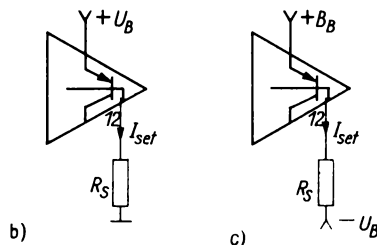
samtverlustleistung von 20 W . Bild 3.49 zeigt den Übersichtsplan der Innenschaltung.

Die oben genannten Schaltungen lassen sich mit integrierten Leistungsoperationsverstärkern viel einfacher realisieren als mit diskreten Bauelementen. Sie arbeiten außerdem wesentlich zuverlässiger, da die Schaltkreise vollständig gegen thermische Überlastung und Kurzschluß geschützt sind.



a)

C nur beim Typ B 176



b)

c)

Bild 3.48. Programmierbarer Kleinleistungs-OPV

a) Übersichtsplan der Innenschaltung (Typ B 177)

b) und c) Einspeisung des Setstromes

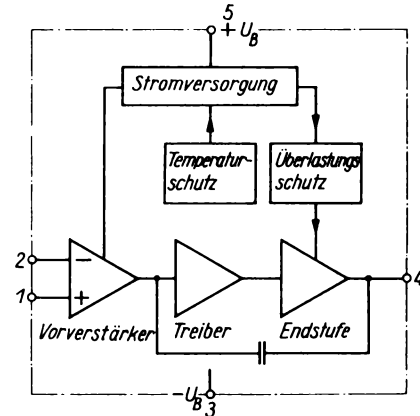


Bild 3.49. Leistungsoperationsverstärker; Übersichtsplan der Innenschaltung (Typ B 165)

3.4.9.

Operationsverstärker mit Stromdifferenzeingang (Norton-Verstärker)

Ein anderer Operationsverstärkertyp weist anstelle des nichtinvertierenden Eingangs eine Stromspiegelschaltung auf, die einen Differenzeingang vortäuscht. Als steuernde Größe wirkt dabei die Differenz der beiden Eingangsströme. Auf Grund der einfachen Schaltung läßt sich dieser Verstärker problemlos in integrierter Technik realisieren. Vorteilhaft ist, daß nur eine Betriebsspannung benötigt wird; trotzdem lassen sich mit ihm die meisten Operationsverstärkerschaltungen verwirklichen.

Bild 3.50 zeigt die vereinfachte Schaltung, ohne die notwendigen Hilfsstufen für die Basisspannungsversorgung für die Transistoren $V5$ und $V7$.

Wesentliches Kennzeichen des Operationsverstärkers mit Stromdifferenzeingang ist die Stromspiegelstufe mit den datengleichen Transistoren $V1$ und $V2$. Diese sichern den notwendigen Eingangsleichtaktbereich. Mit $i_{1P} = B i_B + 2 i_B$ und $i_{C2} = B i_B$ folgt

$$\frac{i_{1P}}{i_{C2}} = 1 + \frac{2}{B} \approx 1, \text{ da } B \gg 1, \text{ also } i_{1P} \approx i_{C2}.$$

Da $i_{1N} = i_{C2} + i_{B3} \approx i_{1P} + i_{B3}$, ergibt sich

$$i_{B3} = i_{1N} - i_{1P}.$$

Der Transistor $V3$ wird also von der Differenz der Eingangsströme in den invertierenden bzw. nichtinvertierenden Eingang gesteuert.

Die Eingänge sind im praktischen Betrieb durch die Basis-Emitter-Spannungen für $V1$ und $V3$ vorgespannt. Zu verstärkende Spannungen müssen deshalb mit Widerständen in Eingangsströme umgewandelt werden.

Der invertierende Verstärker mit $V3$ arbeitet in Emitterschaltung auf die Konstantstromquelle mit $V4$ und $V5$, die als Kollektorwiderstand wirkt. Die Ausgangsstufe mit dem Transistor $V6$ ist eine Kollektorschaltung; als Emitterwiderstand dient die Konstantstromquelle mit $V7$.

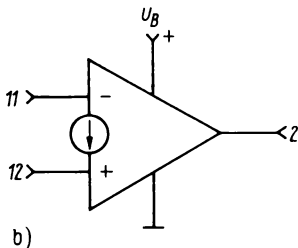
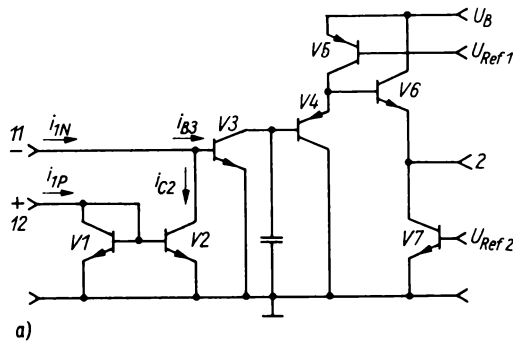


Bild 3.50. Operationsverstärker mit Stromdifferenzeingang

a) Innenschaltung (vereinfacht); b) Schaltzeichen

Zusammenfassung

Ein integrierter Operationsverstärker ist eine standardisierte Universalschaltung für vielseitige Anwendungen in der Elektronik. Mit ihm werden in erster Linie erwünschte Eingangs-Ausgangs-Beziehungen realisiert.

Ein Operationsverstärker soll folgende Eigenschaften aufweisen:

- sehr hohe offene Spannungsverstärkung bis zu beliebig hohen Frequenzen
- sehr hohen Eingangswiderstand
- sehr kleinen Ausgangswiderstand
- geringe Nullpunktfehler und hohe Gleichtaktunterdrückung.

Ein Operationsverstärker benötigt zwei entgegengerichtete Betriebsspannungen mit dem gleichen Absolutwert. In der Praxis kann ein Operationsverstärker in erster Näherung ideal betrachtet werden, wenn er nicht übersteuert wird und Frequenzgang und Offset kompensiert werden.

Ein Operationsverstärker wird stets mit einer äußeren Gegenkopplung betrieben. Diese bestimmt wesentlich die Eigenschaften der Gesamtanordnung.

Man unterscheidet invertierenden und nichtinvertierenden Betrieb. Beim invertierenden Verstärker ist die Ausgangsspannung gegenphasig zur Eingangsspannung, beim nichtinvertierenden dagegen gleichphasig.

Alle integrierten Operationsverstärker haben prinzipiell gleichartige Schaltungskonzepte, sind meist 3stufig ausgeführt und enthalten als Grundsaltungen Differenzverstärker, Konstantstromquelle, Darlington-Schaltung, Koppelschaltungen zur Pegelanpassung und Ausgangsstufe.

Ein anderer Operationsverstärkertyp hat anstelle des nichtinvertierenden Eingangs eine Stromspiegelschaltung. Als steuernde Größe wirkt dabei die Differenz der beiden Eingangsströme. Dieser OPV ist nicht so universell einsetzbar, hat aber den Vorteil, daß nur eine Betriebsspannung benötigt wird.

3.5.

Vorverstärker

3.5.1.

Einstellung des Arbeitspunktes

Wie bereits im Abschn. 3.1. erwähnt, besteht beim Vorverstärker das allgemeine Verstärkerproblem darin, einerseits die Steuerleistung $P_{1\sim}$ ($P_{1\sim} = U_{1\sim} I_{1\sim}$) möglichst klein zu halten, indem man z. B. $I_{1\sim}$ sehr klein hält, und andererseits aus dem Verstärker eine möglichst große Wechselleistung $P_{2\sim} = U_{2\sim} I_{2\sim}$ herauszuholen.

Bei Transistoren benötigt man stets eine Steuerleistung. Der Arbeitspunkt wird dabei so gewählt, daß eine verzerrungsarme Aussteuerung erfolgen kann. Die Basis-Emitter-Spannung U_{BE} muß stets so gerichtet sein, daß der Eingangskreis in Durchlaßrichtung gepolt ist (also negativ bei pnp- und positiv bei npn-Transistoren). Durch folgende Schaltungen kann das realisiert werden.

Einspeisung eines festen Basisstroms (Bild 3.51)

Ein zwischen der Betriebsspannungsquelle und der Basis liegender Vorwiderstand R_1 wird vom Basisstrom durchflossen. Der dabei auftretende Spannungsabfall muß so groß sein, daß die gewünschte Basis-Emitter-Spannung erreicht wird.

Bei gewähltem U_B , U_{BE} und I_B ist

$$U_B = I_B R_1 + U_{BE}$$

$$I_B R_1 = U_B - U_{BE}$$

$$R_1 = \frac{U_B - U_{BE}}{I_B} \quad (3.29a)$$

Mit $|U_B| \gg |U_{BE}|$ folgt näherungsweise aus Gl. (3.29a)

$$R_1 \approx \frac{U_B}{I_B}$$

Da der Spannungsabfall am Vorwiderstand R_1 wesentlich größer als U_{BE} ist, ändert sich der Basisstrom bei temperaturabhängiger Veränderung von U_{BE} nur geringfügig. Der Eingangswiderstand der Schaltung wird im wesentlichen vom Eingangswiderstand des Transistors bestimmt. Sehr nachteilig ist, daß der Arbeitspunkt in Abhängigkeit vom Stromverstärkungsfaktor eingestellt werden muß. Bei Streuung der Transistorparameter macht sich das besonders unangenehm bemerkbar.

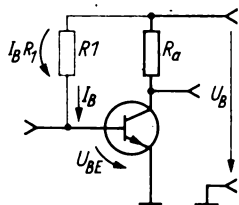


Bild 3.51. Basisvorspannungserzeugung durch einen Vorwiderstand zwischen Basis und Betriebsspannungsquelle

Basisvorspannung mittels Vorwiderstands zwischen Basis und Kollektor

Bezüglich der Arbeitspunktstabilität ist die Schaltung nach Bild 3.52 günstiger als die nach

Bild 3.51. Der Arbeitspunkt wird wieder über den Basisstrom durch den Widerstand R_1 bestimmt, der zwischen Basis und Kollektor liegt. Bei gewähltem U_{CE} , U_{BE} und I_B folgt für den Vorwiderstand R_1

$$R_1 = \frac{U_{CE} - U_{BE}}{I_B} \quad (3.29b)$$

Mit $|U_{CE}| \gg |U_{BE}|$ folgt näherungsweise aus Gl. (3.29b)

$$R_1 \approx \frac{U_{CE}}{I_B}$$

Steigt beispielsweise durch Temperaturänderung der Kollektorstrom an, dann wird der Spannungsabfall an R_1 größer, und die Kollektor-Emitter-Spannung sinkt. Damit wird aber nach Gl. (3.29b) der Basisstrom kleiner und vermindert den ansteigenden Kollektorstrom, so daß letzterer nahezu konstant bleibt. Für eine gute Stabilisierung muß R_1 groß sein.

Es handelt sich hierbei um eine Spannungsgegenkopplung (Parallel-Parallel-Gegenkopplung), die sowohl statisch als auch dynamisch wirkt. Die dynamische Gegenkopplung ist nicht immer erwünscht, da durch sie die Verstärkung und der Eingangswiderstand der Stufe vermindert werden (s. Abschn. 3.3.).

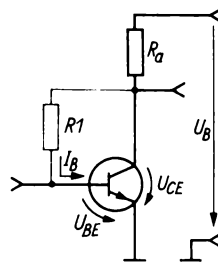


Bild 3.52. Basisvorspannungserzeugung durch einen Vorwiderstand zwischen Basis und Kollektor

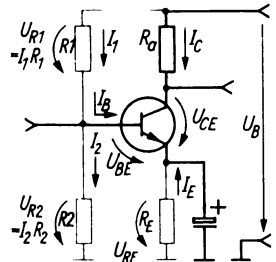


Bild 3.53. Basisvorspannungserzeugung durch einen Spannungsteiler und Gleichstromgegenkopplung

Basisvorspannung durch Spannungsteiler und Gleichstromgegenkopplung (Serien-Serien-Gegenkopplung)

Eine besonders wichtige und wohl die in der Praxis am häufigsten angewendete Schaltung ist die nach Bild 3.53. Sie zeichnet sich durch eine gute Arbeitspunktstabilität aus.

Steigt beispielsweise durch Temperaturerhöhung der Basisstrom und mit ihm der Kollektor- und Emitterstrom, dann wird auch der Spannungsabfall über R_E größer. Das Emitterpotential wird angehoben und damit die Basis-Emitter-Spannung vermindert. Das hat wiederum eine Verkleinerung des Stromes zur Folge, so daß unabhängig von äußeren Einflüssen die Transistorgleichströme nahezu konstant bleiben. Nachteilig dabei ist, daß ein Teil der Betriebsspannung am Emitterwiderstand R_E verlorenggeht. Zur Verkleinerung einer Wechselstromgegenkopplung wird R_E durch einen Kondensator C_E überbrückt.

Der Spannungsteiler wird wie folgt berechnet:

$$R_1 = \frac{U_{R1}}{I_1} = \frac{U_B - (U_{BE} + U_{RE})}{I_1}$$

Mit $I_1 = I_B + I_2$ und $U_{RE} = -I_E R_E = (I_B + I_C) R_E$ folgt

$$R_1 = \frac{U_B - U_{BE} - (I_B + I_C) R_E}{I_B + I_2}$$

$$R_2 = \frac{U_{R2}}{I_2} = \frac{U_{BE} + U_{RE}}{I_2} = \frac{U_{BE} + (I_B + I_C) R_E}{I_2}$$

(3.30a) (3.30b)

Damit einerseits die Betriebsspannungsquelle nicht zu stark belastet und sich andererseits die Spannung an der Basis bei Basisstromänderungen nicht zu stark verändert, ist der durch R_2 fließende Querstrom I_2 so zu wählen, daß er die 2- bis 5fache Stärke des Basisstroms aufweist, d. h. $I_2 \approx (2 \dots 5) I_B$.

Der Emitterwiderstand R_E soll so groß wie möglich sein. Eine praktische Grenze ist aber gegeben durch die Betriebsspannung, von der ein Teil über R_E abfällt.

$$U_B = I_C R_a + U_{CE} + (I_B + I_C) R_E$$

Der Arbeitspunkt im Ausgangskennlinienfeld ist im allgemeinen so zu wählen, daß ein möglichst großer Aussteuerungsbereich vorliegt. Das ist der Fall, wenn die Kollektor-Emitter-Spannung U_{CE} etwa so groß wie der Spannungsabfall am ohmschen Arbeitswiderstand ist.

Beispiel

Ein npn-Transistor soll bei einer Betriebsspannung von $U_B = 12 \text{ V}$ auf einen Arbeitswiderstand $R_a = 1,2 \text{ k}\Omega$ arbeiten. Als Kollektor-Emitter-Spannung sei $U_{CE} = 5 \text{ V}$, als Emitterwiderstand $R_E = 560 \Omega$ gewählt. Dimensionieren Sie die Schaltung!

Gegeben:

$$U_B = 12 \text{ V}$$

$$R_a = 1,2 \text{ k}\Omega$$

$$U_{CE} = 5 \text{ V}$$

$$R_E = 560 \Omega$$

Gesucht:

$$R_1$$

$$R_2$$

Lösung:

Der sich dafür ergebende Kollektorstrom ist etwa

$$I_C \approx \frac{U_B - U_{CE}}{R_a + R_E} = \frac{(12-5) \text{ V}}{1,76 \text{ k}\Omega} = \frac{7}{1,76} \text{ mA} = 4 \text{ mA}$$

Nach dem Kennlinienfeld des verwendeten Transistors fließt dieser Strom bei $I_B \approx 20 \mu\text{A}$ und $U_{BE} \approx 0,68 \text{ V}$.

Der Spannungsteiler R_1/R_2 für die Schaltung nach Bild 3.53 errechnet sich folgendermaßen:

$$R_2 = \frac{U_{BE} + U_{RE}}{I_2}$$

nach Gl. (3.30b) mit $I_2 = 3 I_B = 3 \cdot 20 \mu\text{A} = 60 \mu\text{A}$ und $U_{RE} = (I_B + I_C) R_E = 4,02 \text{ mA} \cdot 560 \Omega \approx 2,25 \text{ V}$

$$R_2 = \frac{(0,68 + 2,25) \text{ V}}{60 \mu\text{A}} = 49 \text{ k}\Omega;$$

gewählt: $R_2 = 47 \text{ k}\Omega$

$$R_1 = \frac{U_B - U_{BE} - (I_B + I_C) R_E}{I_B + I_2}$$

nach Gl. (3.30a)

$$R_1 = \frac{(12 - 0,68 - 2,25) \text{ V}}{(20 + 60) \mu\text{A}} = 113,5 \text{ k}\Omega;$$

gewählt: $R_1 = 110 \text{ k}\Omega$.

Bei mehrstufigen Verstärkern ist zu beachten, daß die einzelnen Stufen durch in die Betriebsspannungsleitung einzuschaltenden RC-Glieder entsprechend Bild 3.54 zu entkoppeln sind.

Das ist deshalb notwendig, damit der Wechselstromkreis für jede einzelne Stufe geschlossen ist und über die Versorgungsspannung keine unerwünschten Rückkopplungen auftreten.

3.5.2.

RC-gekoppelter Verstärker

Die Kopplung mehrerer Verstärkerstufen erfolgt durch Kopplungsvierpole, deren Dimensionierung sich praktisch ausschließlich nach dem zu verstärkenden Frequenzband und der geforderten Verstärkung richtet.

Werden zwei Stufen mit ohmschem Arbeitswiderstand durch einen Kondensator gekoppelt, dann nennt man den Verstärker *Widerstands- oder RC-Verstärker*.

Obwohl eine reine C-Kopplung vorliegt, bezieht man die Art des Arbeitswiderstands in die Benennung mit ein.

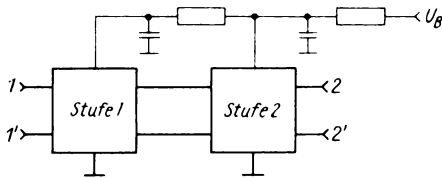


Bild 3.54. Entkopplung der Verstärkerstufen von der Betriebsspannungsquelle

RC-Glied als Kopplungsvierpol

Die Aufgabe eines Kopplungsvierpols besteht darin, einerseits den Eingang des folgenden Verstärkerbauelements gleichspannungsfrei zu machen, andererseits die zu verstärkenden Wechselgrößen zu übertragen. Der Kopplungsvierpol ist also eine elektrische Weiche und enthält stets Blindelemente, die seine Übertragungseigenschaften frequenzabhängig gestalten. Im folgenden wird diese Frequenzabhängigkeit untersucht werden (Bild 3.55).

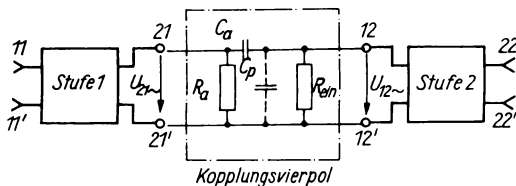


Bild 3.55. Kopplung zweier Verstärkerstufen durch einen RC-Kopplungsvierpol

Entscheidend für den Frequenzgang des Verstärkers ist der Kopplungsvierpol.

C_a ist der Koppelkondensator.

C_p ist die Summe aller Kapazitäten, die als Schaltkapazitäten zwischen den Verbindungsleitungen und als Kapazitäten zwischen den Elektroden der Verstärkerbauelemente am Ausgang der ersten Stufe und Eingang der zweiten Stufe auftreten. C_p ist kein von außen zugeschalteter Kondensator, sondern eine zwischen Elektroden und Leitungen unerwünscht auftretende Schaltungsgröße.

Die Kapazität C_p wird gelegentlich durch einen parallel geschalteten Kondensator vergrößert. Das ist dann nötig, wenn für die Eingangskapazität der zweiten Stufe und die obere Grenzfrequenz exakt reproduzierbare Verhältnisse gefordert werden.

Betrachtet man den Kopplungsvierpol bei sehr tiefen Frequenzen, dann ist wegen ωC_p der durch C_p gebildete Parallelleitwert vernachlässigbar klein.

Dagegen muß der durch C_a gebildete, in Reihe mit R_{ein} liegende Blindwiderstand berücksichtigt werden. Es liegt damit ein aus C_a und R_{ein} bestehender Hochpaß vor. Sein Frequenzverlauf gilt strenggenommen nur für eine konstante Eingangswechselspannung.

Bei hohen Frequenzen kann C_a vernachlässigt werden, da sein Blindwiderstand sehr klein gegenüber R_{ein} ist. Der durch C_p gebildete Blindleitwert muß aber berücksichtigt werden, da er bei hohen Frequenzen in der Größenordnung der Parallelschaltung von $R_a \parallel R_{ein}$ liegt. Man erhält damit einen aus C_p und $R_a \parallel R_{ein}$ gebildeten Tiefpaß. Sein Frequenzverlauf gilt strenggenommen nur für einen konstanten Eingangswechselstrom.

Bei mittleren Frequenzen wird weit oberhalb der Grenzfrequenz des Tiefpasses und weit unterhalb der Grenzfrequenz des Hochpasses gearbeitet, so daß man in erster Näherung ein frequenzunabhängiges Übertragungsverhalten bekommt.

Im Bild 3.56 ist der Frequenzgang des Kopplungsvierpols dargestellt.

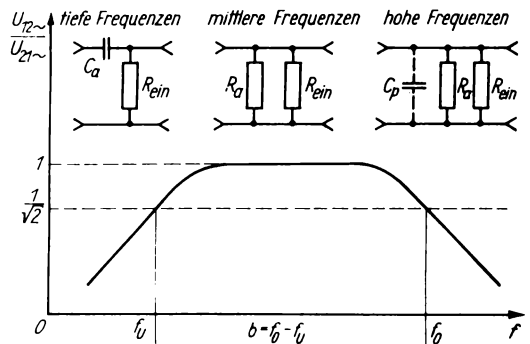


Bild 3.56. Übertragungsverhalten eines RC-Kopplungsvierpols

Bei der unteren und oberen Grenzfrequenz ist der Blindwiderstand gleich dem Wirkwiderstand und damit die Ausgangsspannung des Kopplungsvierpols auf $1/\sqrt{2}$ gegenüber deren Eingangsspannung abgefallen. Man erhält damit leicht die untere und obere Grenzfrequenz.

$$R_{ein} = \frac{1}{\omega_u C_a}$$

$$\text{mit } \omega = 2\pi f$$

$$R_{ein} = \frac{1}{2\pi f_u C_a}$$

$$f_u = \frac{1}{2\pi R_{\text{ein}} C_a} \quad (3.31)$$

$$\frac{1}{R_a} + \frac{1}{R_{\text{ein}}} = \omega_o C_p$$

mit $\omega = 2\pi f$

$$\frac{R_a + R_{\text{ein}}}{R_a R_{\text{ein}}} = 2\pi f_o C_p$$

$$f_o = \frac{1}{2\pi \frac{R_a R_{\text{ein}}}{R_a + R_{\text{ein}}} C_p} \quad (3.32)$$

Frequenzabhängigkeit der Verstärkung

Frequenzabhängigkeit der Emitterkombination

Der Frequenzgang eines Verstärkers wird neben den Übertragungseigenschaften des Kopplungsvierpols auch vom Frequenzverhalten der Emitterkombination beeinflusst.

Bei Transistoren wird zur Arbeitspunktstabilisierung ein Widerstand in die Emitterleitung gelegt und nach Bild 3.53 durch einen Kondensator C_E überbrückt.

Dieser Überbrückungskondensator erfüllt bei hohen Frequenzen seine Aufgabe, indem er die über R_E auftretenden Wechselanteile kurzschließt. Bei tiefen Frequenzen wird sein Blindwert ωC_E kleiner, und es kommt zu einer gewissen Wechselstromgegenkopplung, die die wirksame Steuerwechselgröße vermindert. Dieser Effekt tritt um so stärker in Erscheinung, je niedriger die Frequenz wird, d. h. also, daß die Emitterkombination vor allem die untere Grenzfrequenz beeinflusst.

Bei einer geforderten Grenzfrequenz f_u muß C_E um so größer sein, je kleiner der Generatorwiderstand R_0 ist.

$$C_E \approx \frac{h_{21}}{2\pi f_u (R_0 + h_{11e})} \quad (3.33)$$

Nach Bild 3.58 ist R_0 die Parallelschaltung von R_1 , R_2 und dem Ausgangswiderstand der vorhergehenden Stufe.

Frequenzgang des RC-Verstärkers

Für den Frequenzgang des RC-Verstärkers sind sowohl die Frequenzabhängigkeit der Kopplungsvierpole als auch der Emitterkombination maßgebend. Beide Einflüsse überlagern

sich und ergeben den im Bild 3.57 angegebenen Frequenzverlauf.

Die Frequenzbandbreite ist nach Gl. (3.7)

$$b = f_o - f_u.$$

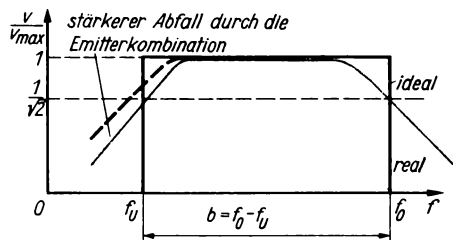


Bild 3.57. Frequenzverlauf eines RC-Verstärkers

Bei *Breitbandverstärkern* ist aber $f_u \ll f_o$, so daß für diese Verstärker näherungsweise $b \approx f_o$ gilt. Mit diesem Ansatz erhält man die interessante Beziehung, daß das Produkt aus Bandbreite mal Verstärkung konstant ist.

Der Nachweis dafür ist leicht zu erbringen. Aus Tafel 3.4 erhält man für $h_{12} = 0$ und $R_a \ll 1/h_{22}$:

$$v_u \approx -\frac{h_{21}}{h_{11}} R_a,$$

wobei R'_a der tatsächliche Arbeitswiderstand bei mittleren Frequenzen ist:

$$R'_a = R_a \parallel R_{\text{ein}} = \frac{R_a R_{\text{ein}}}{R_a + R_{\text{ein}}}.$$

Aus Gl. (3.32) folgt für R'_a an der oberen Grenzfrequenz

$$R'_a = \frac{R_a R_{\text{ein}}}{R_a + R_{\text{ein}}} = \frac{1}{2\pi f_o C_p}.$$

Mit $b \approx f_o$ und $v_u \approx -\frac{h_{21}}{h_{11}} R'_a$ erhält man

$$|v_u| \approx \frac{h_{21}}{h_{11}} \cdot \frac{1}{2\pi b C_p}$$

$$2\pi b |v_u| = \frac{h_{21}}{h_{11}} \cdot \frac{1}{C_p} = \text{konst.},$$

da h_{21}/h_{11} und C_p gegeben sind.

$$\boxed{\text{Bandbreite mal Verstärkung} = \text{konst.}} \quad (3.34)$$

Diese für die Verstärkertechnik wichtige Grundgleichung sagt aus, daß bei einem vorgegebenen Verstärkerbauelement jede Vergrößerung der Bandbreite ein Absinken der Verstärkung zur Folge hat und umgekehrt.

Beim Entwurf eines Verstärkers hat man davon auszugehen, daß sich die Gesamtverstärkung v_{ges} aus dem Produkt der Verstärkungen der einzelnen Stufen ergibt.

$$v_{\text{ges}} = v_1 v_2 v_3 \dots \quad (3.35)$$

Damit ist aber der gesamte Frequenzgang ebenfalls das Produkt der einzelnen Stufenfrequenzgänge.

$$\left(\frac{v}{v_{\text{max}}}\right)_{\text{ges}} = \left(\frac{v}{v_{\text{max}}}\right)_1 \left(\frac{v}{v_{\text{max}}}\right)_2 \left(\frac{v}{v_{\text{max}}}\right)_3 \quad (3.36)$$

Der Verstärkungsabfall bei tiefen und hohen Frequenzen ist dadurch bei mehrstufigen Verstärkern viel steiler. Liegt beispielsweise ein 2stufiger Verstärker mit völlig gleichen Stufen vor und der Verstärkungsabfall einer Stufe bei der Frequenz $f_{u,2}$ und $f_{o,2}$ beträgt $v/v_{\text{max}} = 1/\sqrt{2} = 0,707$, dann ergibt sich bei den gleichen Frequenzen ein Gesamtabfall von

$$\left(\frac{v}{v_{\text{max}}}\right)_{\text{ges}} = \frac{1}{\sqrt{2}} \cdot \frac{1}{\sqrt{2}} = 0,5.$$

Wird dagegen bei den Grenzfrequenzen f_u und f_o ein Gesamtabfall von $\left(\frac{v}{v_{\text{max}}}\right)_{\text{ges}} = \frac{1}{\sqrt{2}}$ gefordert, dann darf jede Stufe bei diesen Frequenzen nur einen Abfall von

$$\left(\frac{v}{v_{\text{max}}}\right)_{1,2} = \sqrt[2]{\frac{1}{\sqrt{2}}} = \sqrt[2]{0,707} = 0,878 \approx 0,88$$

aufweisen.

Schaltung des RC-Verstärkers

Im Bild 3.58 wird die praktische Schaltung eines 2stufigen RC-Vorverstärkers wiedergegeben.

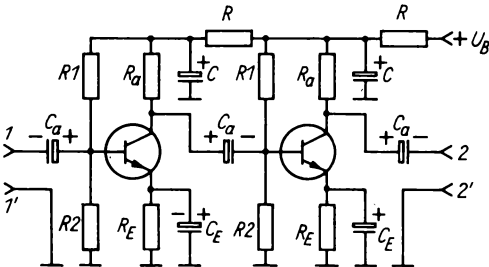


Bild 3.58. Zweistufiger RC-Vorverstärker

3.5.3.

Bandfiltergekoppelter Verstärker

Bei diesem Verstärker werden als Kopplungs-vierpole Bandfilter eingesetzt. Diese sind stark frequenzabhängige Netzwerke und übertragen nur ein bestimmtes schmales Frequenzband. Aus diesem Grund bezeichnet man bandfiltergekoppelte Verstärker auch als *Schmalband- oder Selektivverstärker*.

Sie werden hauptsächlich in der Hochfrequenz-technik benötigt, z. B. als Zwischenfrequenzverstärker (ZF-Verstärker) im Rundfunk- und Fernsehempfänger. Bandfilter können sehr verschieden aufgebaut sein. Neben reinen RC-Filtern gibt es LC-Filter, die sowohl aus einem als auch aus mehreren Kreisen bestehen können, und mechanische Bandfilter. Bei ZF-Verstärkern benutzt man meistens induktiv oder kapazitiv gekoppelte zweikreisige LC-Bandfilter entsprechend Bild 3.59 und in zunehmendem Maß auch mechanische Bandfilter. Bei ersteren kann die Durchlaßkurve durch die Größe der Kreiskopplung variiert werden. Im Bild 3.60 gilt der Verlauf a für sehr lose („unterkritische“), b für mittlere („kritische“) und c für feste („überkritische“) Kopplung. Es kann damit ohne Verminderung der „Flankensteilheit“ der Durchlaßkurve die Bandbreite vergrößert werden.

Zu beachten ist bei diesen Verstärkern die Rückwirkungsgröße der Verstärkerbauelemente.

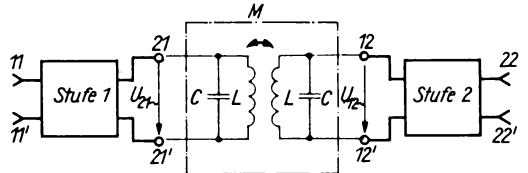


Bild 3.59. Kopplung zweier Verstärkerstufen durch ein induktiv gekoppeltes zweikreisiges LC-Bandfilter

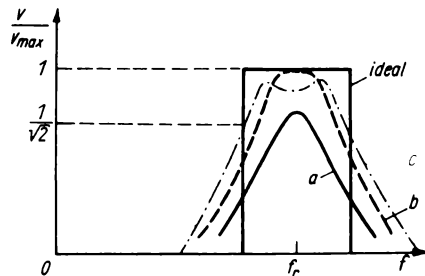


Bild 3.60. Durchlaßkurven eines bandfiltergekoppelten Verstärkers bei verschieden fester Kopplung

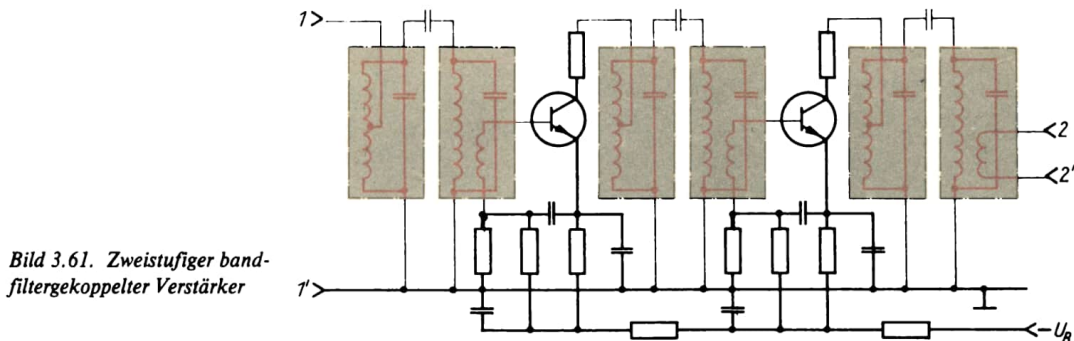


Bild 3.61. Zweistufiger bandfiltergekoppelter Verstärker

Da an den Ein- und Ausgängen Schwingkreise mit der gleichen Resonanzfrequenz liegen, kann es durch Mitkopplung zur Selbsterregung der Schaltungsanordnung kommen. Diese Gefahr wird um so größer, je höher die Frequenz der zu verstärkenden Wechselgröße ist. Durch entsprechend dimensionierte Neutralisierungsglieder muß diese unerwünschte Erscheinung beseitigt werden.

Als Beispiel wird im Bild 3.61 eine auf Feinheiten verzichtende Schaltung eines Selektivverstärkers gezeigt. Die Bandfilter sind kapazitiv gekoppelte KleinfILTER. Eine kleine Koppelspule führt die Ausgangsgröße an die jeweils nächste Stufe.

Dem verstimmenden Einfluß der Transistoren bei Betriebsspannungs- oder Aussteuerungsänderungen wird durch die Kollektorwiderstände und Anzapfung der Schwingkreise entgegengewirkt. Die in der Betriebsspannungszuführung liegenden RC-Glieder verhindern eine unerwünschte Verkopplung der einzelnen Stufen.

3.5.4.

Direktgekoppelter Verstärker

Alle bisher behandelten Verstärker versagen, wenn es darauf ankommt, Gleichspannungen oder besser Spannungsschwankungen sehr niedriger Frequenz oder niederfrequente Impulse zu verstärken. Bei ihnen ist wegen der frequenzabhängigen Kopplungsvierpole eine in verschiedenen Fällen notwendige untere Grenzfrequenz $f_u = 0$ Hz (z. B. bei Verstärkern in der Meß- und Regelungstechnik) nicht erreichbar. Man muß dann eine direkte Kopplung anwenden, indem man den Ausgang einer Verstärkerstufe mit dem Eingang der folgenden galvanisch verbindet.

Einen derartig ausgeführten 2stufigen Verstärker zeigt Bild 3.62.

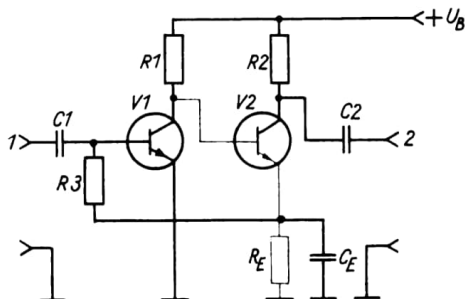


Bild 3.62. Zweistufiger direktgekoppelter Verstärker

Die beiden Transistoren V1 und V2 arbeiten für Wechselspannungen ohne Gegenkopplung in Emitterschaltung. Durch den Widerstand R_E wird das Emittential des Transistors V2 angehoben, so daß eine direkte Kopplung vom Kollektor des Transistors V1 auf die Basis des Transistors V2 möglich ist. Gleichzeitig erfolgt durch R_E eine Gleichspannungsgegenkopplung, die sich über R_3 auch voll auf den Transistor V1 auswirkt. Die Arbeitspunkte sind damit stabilisiert.

Bei direktgekoppelten Verstärkern ist auch eine Wechselspannungsgegenkopplung über mehrere Stufen leichter realisierbar als bei RC-gekoppelten, weil dabei keine Verstärkungs- und Phasenbeziehungen an der unteren Grenzfrequenz realisiert werden müssen.

Bild 3.63 zeigt einen derartigen gegengekoppelten Verstärker.

Dem 2stufigen Verstärker ist eine Kollektorschaltung durch direkte Kopplung angeschlossen. Durch die Widerstände R_5 und R_6 erfolgt eine Spannungsgegenkopplung. Die Arbeitspunktstabilisierung erfolgt genauso wie in der Schaltung nach Bild 3.62. Da R_5 sehr niederohmig sein kann und R_6 im Interesse der Verstärkung im allgemeinen groß gegenüber R_5 ge-

wählt wird, ist $V1$ praktisch nicht stromgegegenggekoppelt.

Die Vorteile dieses gegengekoppelten gegenüber dem nichtgegengekoppelten Verstärker liegen darin, daß der Verstärkungsfaktor von den Transistordaten weitgehend unabhängig ist, Ausgangswiderstand und Klirrfaktor verkleinert, Eingangswiderstand und Bandbreite vergrößert werden.

Bei Verwendung von Komplementärtransistoren lassen sich direktgekoppelte Verstärker noch einfacher aufbauen, da das Aufstocken der Betriebsspannung von Stufe zu Stufe entfallen kann.

Bild 3.64 zeigt einen derartigen 2stufigen Verstärker.

Beide Stufen sind durch R_{E1} bzw. R_{E2} stromgegegenggekoppelt. Da sich damit nur kleine Verstärkungen erreichen lassen, wird diese Schaltung selten als Spannungsverstärker, sondern vor allem als Impedanzwandler eingesetzt. Eine sehr starke Vereinfachung ergibt sich, wenn der Emittter des Transistors $V1$ zum Kollektor des Transistors $V2$ geführt und der Widerstand R_{E2} Null gemacht wird (Bild 3.65).

Die so gewonnene Anordnung ist eine Komplementär-Darlington-Schaltung, die wie ein Emittterfolger mit einem npn-Transistor arbeitet, eine Spannungsverstärkung $v_u = 1$ aufweist und einen sehr großen Eingangs- und sehr kleinen Ausgangswiderstand hat.

Vorverstärker werden auch als integrierte Schaltkreise gefertigt. Diese IS haben neben einer hohen Zuverlässigkeit sehr kleine Abmessungen, geringe Masse und eine hohe mechanische Stabilität. Bei diesen Schaltkreisen werden die einzelnen Stufen stets direkt gekoppelt. Durch äußere Beschaltung mit einigen wenigen diskreten Bauelementen lassen sich damit komplette Verstärker mit solchen Kenndaten aufbauen, die den jeweiligen praktischen Erfordernissen entsprechen. Als Beispiel sei im Bild 3.66 die Schaltung eines NF-Vorverstärkers mit dem Schaltkreis MAA 145 des Werkes Tesla Roznov (ČSSR) wiedergegeben.

Direktgekoppelte Verstärker können auch als Selektivverstärker aufgebaut werden, indem man eine frequenzabhängige Gegenkopplung einsetzt. Allerdings erhält man dabei keine solche rechteckige Durchlaßkurve wie bei einem bandfiltergegekoppelten Verstärker.

Als Gegenkopplungsnetzwerk kann beispielsweise ein Doppel-T-Glied verwendet werden (Bild 3.67). Dieses verhält sich so, daß tiefe Fre-

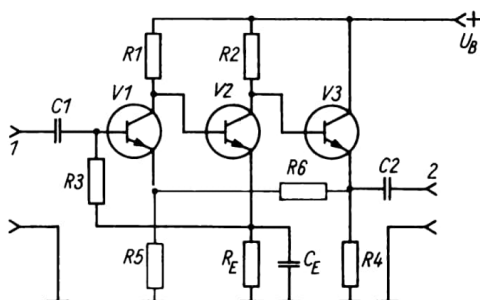


Bild 3.63. Direktgekoppelter Verstärker mit Gegenkopplung

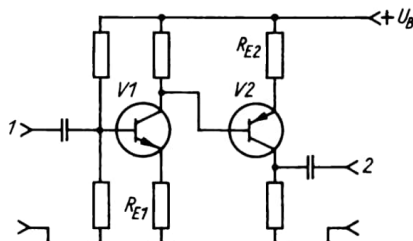


Bild 3.64. Direktgekoppelter Verstärker mit Komplementärtransistoren

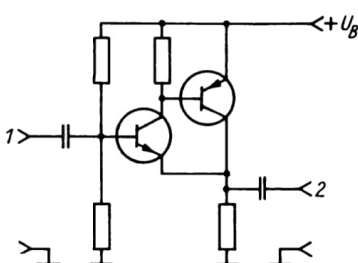


Bild 3.65. Direktgekoppelter Verstärker als Impedanzwandler

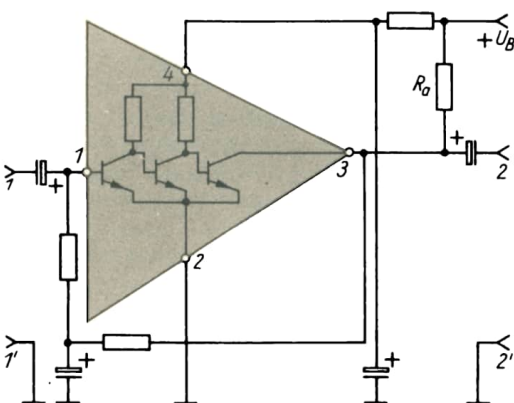


Bild 3.66. Dreistufiger Verstärker mit dem integrierten Schaltkreis MAA 145

quenzen über die beiden Widerstände R , hohe dagegen über die beiden Kondensatoren C voll übertragen werden. Bei der Frequenz

$$f_r = \frac{1}{2\pi RC}$$

ist die Ausgangsspannung Null.

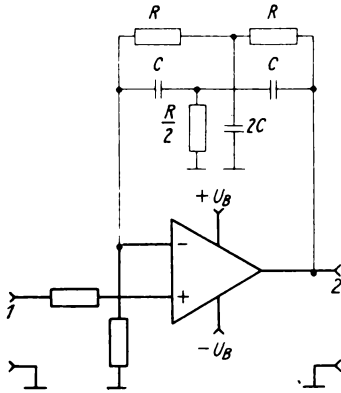


Bild 3.67. Selektivverstärker durch frequenzabhängige Gegenkopplung

Wird nun das Doppel-T-Glied als Gegenkopplungsnetzwerk in einen Operationsverstärker eingesetzt, dann nimmt mit steigender Eingangsfrequenz die an dem invertierenden Eingang zurückgekoppelte Spannung langsam ab. Bei der Nullfrequenz f_r wird die Gegenkopplung unwirksam, und der OPV arbeitet nahezu im Leerlauf. Wird die Eingangsfrequenz über f_r weiter vergrößert, dann erhöht sich die Gegenkopplung wieder und die Verstärkung sinkt ab. Bild 3.68 zeigt die Durchlaßkurve dieses Selektivverstärkers.

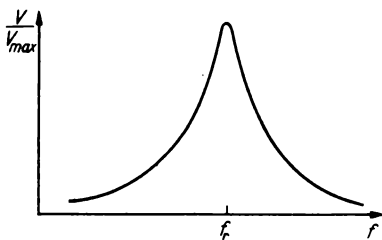


Bild 3.68. Durchlaßkurve des Selektivverstärkers nach Bild 3.67

3.5.5.

Optoelektronisch gekoppelter Verstärker

Zur galvanischen Trennung von Stromkreisen mit sehr hohen Potentialdifferenzen verwendet man als Kopplungsvierpol zwischen den Verstärkerstufen optoelektronische Koppler (Bild 3.69).

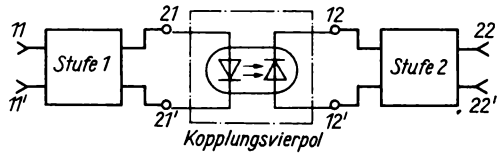


Bild 3.69. Kopplung zweier Verstärkerstufen durch einen optoelektronischen Koppler

Ein optoelektronischer Koppler hat in einem Gehäuse einen Lichtsender und einen Lichtempfänger. Als Sender dient eine lichtemittierende Diode, die bei Stromdurchgang infrarotes oder sichtbares Licht abstrahlt. Als Empfänger verwendet man Fotowiderstände, Fototransistoren, Fotodioden und Fotodioden mit integriertem Verstärker. In analogen Schaltungen bevorzugt man Fotodiodenkoppler. Gegenüber Kopplern mit Fotowiderstand oder Fototransistor weisen diese geringere Verzögerungszeiten, größere Bandbreiten und bessere Linearität der Übertragungsfunktion auf.

Optoelektronisch gekoppelte Verstärker werden vorwiegend in Anlagen der Steuerungs- und Regelungstechnik eingesetzt, gelegentlich aber auch in Geräten der Unterhaltungselektronik. Bild 3.70 zeigt die Schaltung eines Verstärkers mit optoelektronischem Koppler. Der Transistor $V1$ sorgt für die Aussteuerung und den Ruhestrom der Lichtemitterdiode. Der nachfolgende Verstärker mit den Transistoren $V2$ und $V3$ ist gegengekoppelt.

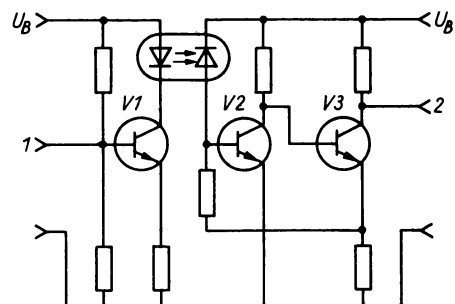


Bild 3.70. Optoelektronisch gekoppelter Verstärker

Zusammenfassung

Vorverstärker haben im allgemeinen eine hohe Spannungsverstärkung. Sie arbeiten verzerungsarm.

Bei mehrstufigen Verstärkern werden die einzelnen Stufen durch Kopplungsvierpole verbunden. Je nach Art der Koppellemente unterscheidet man u. a. RC-Verstärker, bandfiltergekoppelte, direktgekoppelte und optoelektronisch gekoppelte Verstärker. Der Frequenzgang der Kopplungsvierpole ist bestimmend für den Gesamtfrequenzgang des Verstärkers. Durch geeignete Dimensionierung der Koppellemente und zusätzliche Schaltungsmaßnahmen (Gegenkopplung) kann der Frequenzgang den praktischen Bedürfnissen entsprechend gestaltet werden. Jedoch ist zu beachten, daß wegen der Beziehung

$$\text{Bandbreite mal Verstärkung} = \text{konst.}$$

jede Vergrößerung der Frequenzbandbreite einen Verstärkungsrückgang zur Folge hat.

RC-, direktgekoppelte und optoelektronisch gekoppelte Verstärker sind meist Breitbandverstärker. Durch eine frequenzabhängige Gegenkopplung können sie aber auch selektiv wirken.

Bandfiltergekoppelte Verstärker sind immer Selektivverstärker. Bei diesen ist die Gefahr einer Mitkopplung über die Rückwirkungsgröße der Transistoren besonders bei höheren Frequenzen groß und muß durch geeignete Schaltungsmaßnahmen neutralisiert werden.

Integrierte Vorverstärker sind direktgekoppelte Verstärker und haben infolge ihrer großen Bauelementedichte sehr kleine Abmessungen. Weitere Vorzüge sind ihre hohe Zuverlässigkeit, ihre geringe Masse und ihre große mechanische Stabilität. Sie werden deshalb bevorzugt eingesetzt. Mit optoelektronischen Kopplern lassen sich Verstärkerstufen mit unterschiedlichem Bezugspotential verbinden, wobei zwischen den Bezugspunkten ein sehr großer Potentialunterschied existieren kann.

3.6.

Endverstärker

3.6.1.

Allgemeines

Der Endverstärker soll eine möglichst große verzerrungsarme Wechselleistung nach außen abgeben. Aus diesem Grund können die bei den

Vorverstärkern gemachten Überlegungen nicht einfach übernommen werden, da der Verstärkungsgrad gegenüber einer hohen Ausgangsleistung bei kleinem Klirrfaktor eine untergeordnete Rolle spielt.

Man unterscheidet die Endverstärker nach der Lage des Arbeitspunktes und definiert die Betriebsarten A, AB, B und C (Bild 3.71). Beim A-Betrieb wird symmetrisch um den Arbeitspunkt angesteuert, der etwa auf der Mitte des „geradlinigen“ Teiles der Steuerkennlinie liegt. A-Betrieb ist notwendig bei Eintaktendstufen.

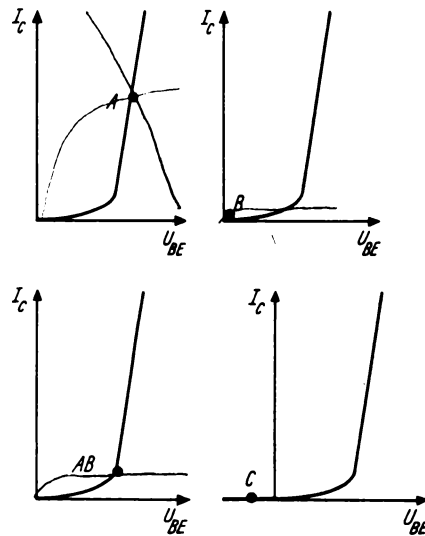


Bild 3.71. Lage des Arbeitspunktes bei den verschiedenen Verstärkertypen

Beim B-Betrieb liegt der Arbeitspunkt bei sehr kleinem Ausgangsruhestrom, so daß die Kollektorspannung im Arbeitspunkt etwa gleich der Betriebsspannung ist.

Beim AB-Betrieb wird, abhängig von der jeweiligen Größe des Eingangssignals, der Arbeitspunkt zwischen A- und B-Betrieb verschoben. Damit werden die Eigenschaften eines A- mit denen eines B-Verstärkers vereinigt.

Beim C-Betrieb liegt der Arbeitspunkt so weit im Sperrbereich des Verstärkerbauelements, daß nur so lange ein Strom fließt, wie die Eingangsschwingung über den Kennlinienfußpunkt hinausschwingt. Es entstehen also nur kurze Stromspitzen mit längeren stromlosen Pausen. C-Betrieb ist auf Resonanzverstärkung beschränkt und findet Anwendung im Senderverstärker.

Da die Steuerkennlinien der Transistoren nicht-

linear sind, wird bei größerer Aussteuerung die Ausgangsgröße verzerrt. Ein Maß für die nichtlineare Verzerrung ist der nach Gl. (3.8) definierte Klirrfaktor. Steuert man beispielsweise einen Transistor sinusförmig bei einem im gekrümmten Kennlinienfeld liegenden Arbeitspunkt aus, dann verläuft der Ausgangswechselstrom entsprechend Bild 3.72 nicht mehr sinusförmig.

Diese verzerrte Sinuskurve kann man mathematisch (Fourier-Analyse) in reine Sinuskurven zerlegen, die die ein-, zwei- bzw. n-fache Frequenz der angelegten Steuerwechselspannung aufweisen.

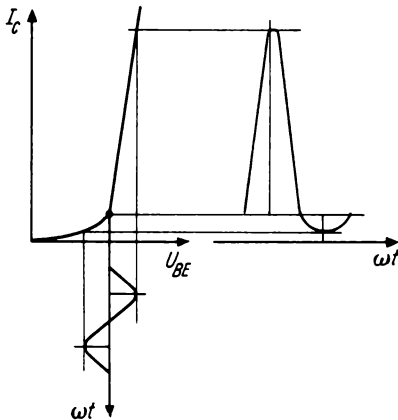


Bild 3.72. Entstehung nichtlinearer Verzerrungen

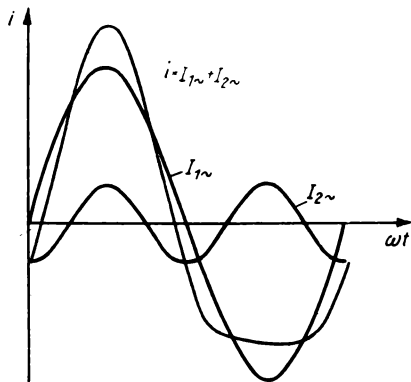


Bild 3.73. Verzerrte Sinusschwingung, bestehend aus erster und zweiter Harmonischen

Man bezeichnet diese Frequenzen als Oberwellen bzw. Harmonische, wobei die Grundfrequenz die erste Harmonische, die erste Oberwelle die zweite Harmonische usw. ist. Andererseits lassen sich natürlich auch reine Sinus-

schwingungen zusammensetzen, indem man deren Momentanwerte addiert. Im Bild 3.73 sind die ersten und zweiten Harmonischen zu einer neuen, verzerrten Schwingung zusammengesetzt.

Oft spielt aber nicht nur *eine* Oberwelle eine Rolle, sondern theoretisch unendlich viele. Die Zahl der zu berücksichtigenden Oberwellen hängt von der gegebenen Kennlinienform und dem Aussteuerungsgrad ab. Im dargestellten Beispiel mit nur einer Oberwelle $I_{2\sim}$, deren Effektivwert $I_{2\sim} = \frac{1}{4} I_{1\sim}$ gegenüber dem der ersten

Harmonischen $I_{1\sim}$ beträgt, ist der Klirrfaktor k entsprechend Gl. (3.8)

$$k = \sqrt{\frac{I_{2\sim}^2}{I_{1\sim}^2 + I_{2\sim}^2}} = \sqrt{\frac{\frac{1}{16}}{1 + \frac{1}{16}}} = \sqrt{\frac{1}{17}} = 24,25 \, \%$$

Bei Musikübertragungen sollte der Klirrfaktor 5 % nicht übersteigen. An hochwertige Anlagen werden noch wesentlich strengere Anforderungen gestellt. Liegen am Verstärkereingang gleichzeitig Signale mit mehreren Frequenzen an, so treten durch nichtlineare Verzerrungen außerdem Mischwellen auf. Diese Mischwellen sind unharmonisch und beeinflussen das Klangbild noch mehr. Man nennt sie auch „Kombinationstöne“.

3.6.2.

Eintaktverstärker

Im Bild 3.74 ist die grundsätzliche Schaltung eines Eintakt-A-Endverstärkers in Emitterschaltung wiedergegeben. Der meist niederohmige Lautsprecherwiderstand R_L wird durch den Ausgangsübertrager hochtransformiert auf den für

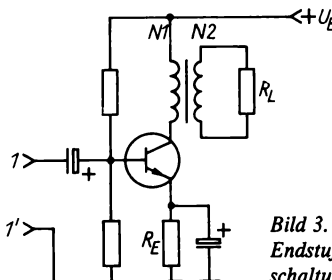


Bild 3.74. Eintakt-A-Endstufe in Emitterschaltung

den Transistor günstigen Wert

$$R_a = \ddot{u}^2 R_L$$

wobei das Übersetzungsverhältnis

$$\ddot{u} = \frac{N_1}{N_2} > 1.$$

Bei der Dimensionierung der Schaltung geht man zweckmäßigerweise vom Ausgangskennlinienfeld Bild 3.75 aus.

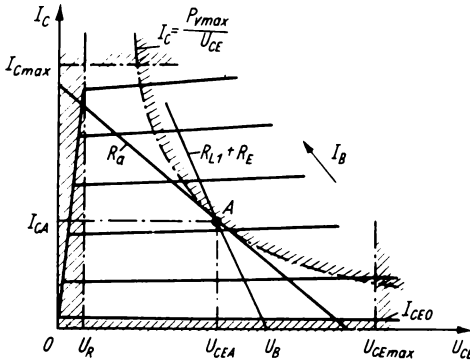


Bild 3.75. Festlegung des Arbeitspunktes und des Arbeitswiderstandes im Ausgangskennlinienfeld

Die Widerstandsgerade $R_{L1} + R_E$ beginnt bei der meist vorgegebenen Betriebsspannung. Auf den Berührungspunkt dieser Geraden mit der Verlustleistungshyperbel $I_C = P_{V\max}/U_{CE}$ legt man bei geforderter maximaler Ausgangsleistung den Arbeitspunkt mit den Werten I_{CA} und U_{CEA} . Die Arbeitswiderstandsgerade R_a tangiert nun die Hyperbel, wobei zu beachten ist, daß die vom Transistorhersteller angegebenen Maximalwerte für Kollektorstrom und -spannung nicht überschritten werden, da anderenfalls der Arbeitspunkt unterhalb der Verlustleistungshyperbel gelegt werden muß. Auf der für einen phasenreinen Widerstand geltenden Arbeitsgeraden R_a erfolgt die Aussteuerung, die durch die Restspannung und den Kollektorreststrom begrenzt wird. Für überschlagsmäßige Berechnungen können die beiden letzten Größen wegen ihrer Kleinheit meist vernachlässigt werden. Da die Arbeitswiderstandsgerade R_a die Tangente an der Verlustleistungshyperbel sein soll, wird unter obenstehender Vernachlässigung

$$R_a = \frac{U_{CEA}}{I_{CA}} = \frac{U_{CEA}^2}{P_{V\max}} \quad (3.38)$$

Die aus den Grundlagen bekannte Anpassungsbeziehung $R_a = R_{aus}$ für maximale Leistungsabgabe kann nicht verwirklicht werden, da dann eine hohe Betriebsspannung erforderlich wäre und der Transistor einen unreal hohen Grenzwert für die maximale Kollektorspannung haben müßte.

Ein Endverstärker läßt sich auch als Emitterfolger entsprechend Bild 3.76 aufbauen.

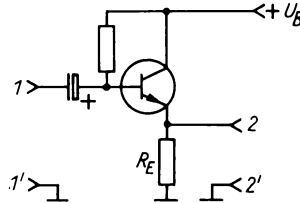


Bild 3.76. Eintakt-A-Endstufe in Kollektorschaltung

Maximale Leistung wird dann abgegeben, wenn Leistungsanpassung vorliegt, also der anzuschließende Verbraucher $R_a = R_E$ wird. Wegen der relativ niedrigen Leistungsverstärkung des Emitterfolgers ist aber der Wirkungsgrad der Schaltung mit etwa 6 % sehr gering. Diese Schaltung wird nur noch selten angewendet.

3.6.3.

Gegentaktverstärker

Wirkungsweise und Arten von Gegentaktverstärkern

Durch Gegentaktschaltung zweier Endtransistoren läßt sich die Ausgangsleistung vergrößern, und die nichtlinearen Verzerrungen gegenüber einer Eintaktschaltung werden erheblich vermindert.

Es gibt vier prinzipielle Schaltungsmöglichkeiten für Gegentaktverstärker (Bild 3.77):

- Parallel-Gegentaktschaltung
- Serien-Gegentaktschaltung

mit jeweils Transistoren gleichen oder unterschiedlichen Leitungstyps.

Bei den Parallel-Gegentaktschaltungen liegen die Transistoren gleichstrommäßig parallel, bei den Serien-Gegentaktschaltungen in Reihe. Werden Transistoren gleichen Leitungstyps eingesetzt, sind zwei um 180° phasenverschobene Steuersignale erforderlich; Schaltungen mit Transistoren unterschiedlichen Leitungstyps müssen gleichphasig, also unmittelbar mit ein und demselben Signal angesteuert werden.

Bei Parallel-Gegentaktschaltungen müssen die beiden Arbeitswiderstandshälften verkettet wer-

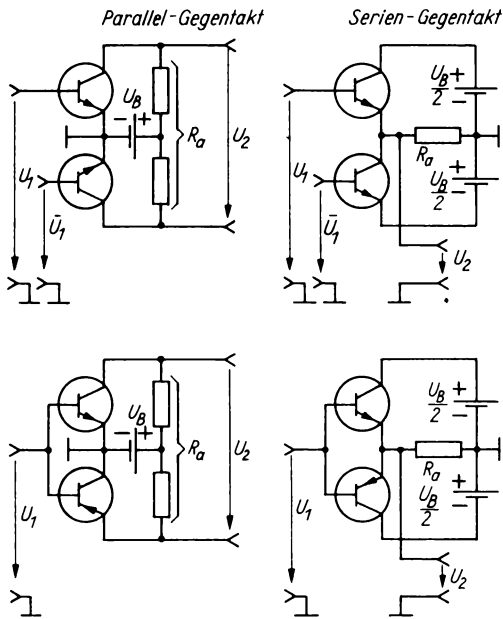


Bild 3.77. Parallel- und Serien-Gegentaktschaltungen mit Transistoren gleichen und unterschiedlichen Leitungstyps

den. Man benötigt dazu einen Transformator, dessen Primärwicklung eine Mittelanzapfung hat. Über diese wird den Transistoren die Betriebsspannung zugeführt. An der Sekundärwicklung ist der Arbeitswiderstand angeschaltet. Da derartige Transformatoren im Vergleich zu

anderen Bauelementen der Schaltung teuer sind, ein großes Volumen haben, die Bandbreite des Verstärkers begrenzen und nicht integrierbar sind, werden Parallel-Gegentaktverstärker nur noch in wenigen Sonderfällen aufgebaut. Bei Serien-Gegentaktschaltungen ist keine transformatorische Ankopplung des Arbeitswiderstands erforderlich. Das ist der entscheidende Grund, warum in der modernen Schaltungstechnik dieses Schaltungsprinzip dominiert. Die folgenden Ausführungen beschränken sich deshalb auf Serien-Gegentaktverstärker. Bezüglich der Betriebsart ergeben sich folgende Verhältnisse (s. auch Tafel 3.6).

Beim *A-Betrieb* ist die von der Schaltung aufgenommene Gesamtleistung unabhängig von der Aussteuerung konstant; der Wirkungsgrad ist maximal 50 %. Vorteilhaft ist auch, daß sich alle geradzahigen Harmonischen gegenseitig aufheben. Das bedeutet eine sehr erhebliche Verminderung der Verzerrungen, da die zweite Harmonische (erste Oberwelle) im allgemeinen den Hauptanteil an den Verzerrungen bildet. Beim *B-Betrieb* ist die aufgenommene Gesamtleistung aussteuerungsabhängig; der Wirkungsgrad ist mit etwa 78 % sehr hoch. Wegen des niedrigen Ruhestroms eignet sich der B-Betrieb vor allem für batteriebetriebene Transistorgehäte. Nachteilig ist, daß bei kleinen Aussteuerungen relativ große Übernahmeverzerrungen auftreten.

Tafel 3.6. Endverstärker

Verstärkertyp	Arbeitspunkt	Bevorzugte Anwendung
Eintaktverstärker		
A-Betrieb	liegt etwa auf der Mitte des „geradlinigen“ Teiles der Steuerkennlinie	als Endverstärker geringerer Leistung
Gegentaktverstärker		
A-Betrieb	liegt auf etwa der Mitte des „geradlinigen“ Teiles der Steuerkennlinie beider aktiven Verstärkerbauelemente	als Endverstärker mittlerer Leistung und kleinem Klirrfaktor bei geringer Aussteuerung
AB-Betrieb	liegt zwischen A- und B-Betrieb und ist stark aussteuerungsabhängig	als Endverstärker relativ großer Leistung und kleinem Klirrfaktor bei geringer und großer Aussteuerung
B-Betrieb	liegt im Fußpunkt der Steuerkennlinie beider aktiven Verstärkerbauelemente (etwa 5 % des Maximalstroms)	als Endverstärker relativ großer Leistung und kleinem Klirrfaktor bei relativ großer Aussteuerung. Wegen des geringen Ruhestrombedarfs Anwendung in portablen, batteriebetriebenen Geräten
C-Betrieb	liegt weit im Sperrbereich der aktiven Verstärkerbauelemente	als Resonanzverstärker hoher Leistung und großen Wirkungsgrads im Senderbetrieb

Der *Gegentakt-AB-Verstärker* nimmt hinsichtlich erreichbarer Ausgangsleistung und Klirrfaktor eine Mittelstellung ein, da sein Arbeitspunkt zwischen A- und B-Betrieb liegt. Bei kleinen Aussteuerungen arbeitet der Gegentakt-AB-Verstärker nahezu im A-Betrieb; bei großen dagegen wirkt er mehr und mehr als B-Verstärker. Endverstärker sind Schaltungen, bei denen eine hohe Ausgangsleistung im Vordergrund steht. Die Spannungsverstärkung spielt dabei eine untergeordnete Rolle und liegt in der Größenordnung $v_u = 1$. Da zur Aussteuerung der Endverstärker aber eine Steuerleistung erforderlich ist, die vom Vorverstärker nicht immer aufgebracht werden kann, muß diese durch eine als *Treiberstufe* bezeichnete Vorstufe geliefert werden. Die Treiberstufe nimmt eine Mittelstellung zwischen Vor- und Endverstärker ein. Ihre Dimensionierung richtet sich nach der aufzubringenden Steuerleistung für die Endstufe.

Gegentaktverstärker mit Transistoren gleichen Leitungstyps

Bei Verwendung von Transistoren gleichen Leitungstyps sind die im Bild 3.78 dargestellten Varianten am einfachsten zu übersehen.

Von den beiden Endstufentransistoren wird *V1* in Kollektor- und *V2* in Emitterschaltung betrieben.

V1 und *V2* werden mit zwei gleich großen, aber gegenphasigen Spannungen angesteuert und daher abwechselnd leitend. Ist beispielsweise die

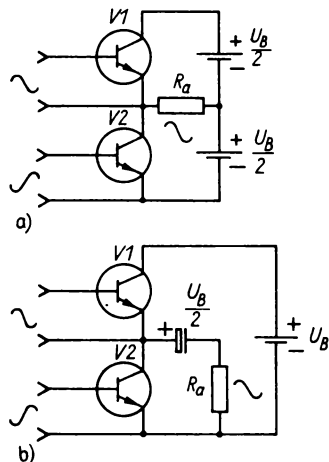


Bild 3.78. Prinzip des Serien-Gegentaktverstärkers mit npn-Transistoren

a) mit zwei Betriebsspannungen; b) mit einer Betriebsspannung

Basis des oberen Transistors positiv, gleichzeitig die des unteren negativ, dann liefert nur der obere Transistor *V1* Strom. Bei Polaritätswechsel sperrt *V1*, und der untere Transistor *V2* wird leitend. Im Arbeitswiderstand R_a werden die beiden übertragenen Halbwellen also wieder zu einer vollen Schwingung zusammengesetzt. Schaltungsvariante a benötigt eine Betriebsspannungsquelle mit Mittelanzapfung. An jedem Transistor liegt dann die Spannung $U_B/2$. Steht keine Spannungsquelle mit Mittelanzapfung zur Verfügung, muß die Schaltungsvariante b gewählt werden. An jedem Transistor liegt dann ebenfalls die Spannung $U_B/2$. Nachteilig ist allerdings, daß in Reihe mit dem Arbeitswiderstand (z. B. Lautsprecher) im Interesse einer tiefen unteren Grenzfrequenz ein sehr großer Kondensator von 500...2000 μF vorzusehen ist.

Gegentaktverstärker mit Transistoren des gleichen Leitungstyps müssen mit zwei gleich großen, aber gegenphasigen Wechselspannungen angesteuert werden. Diese werden von einer *Phasenumkehrstufe* bei anliegendem Eingangssignal zur Verfügung gestellt.

Sehr einfach läßt sich eine Phasenumkehr durch einen Transformator mit zwei Sekundärwicklungen (bzw. Sekundärwicklung mit Mittelabgriff) erreichen. Anstelle des komplizierten Gegentakt-Eingangstransformators mit den bekannten Nachteilen wendet man bevorzugt Transistorschaltungen zur Phasenumkehr an. Eine prinzipielle Lösung wird im Bild 3.79 gezeigt.

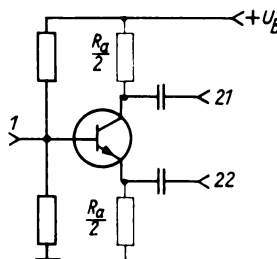


Bild 3.79
Phasenumkehrstufe

Im Ausgangskreis des Transistors liegen zwei gleich große Arbeitswiderstände $R_a/2$, über die bei Aussteuerung auch zwei gleich große, aber gegenphasige Spannungen abfallen müssen. Die Eingangsspannung zwischen Basis und Masse ist um U_{BE} größer als die über den unteren Arbeitswiderstand auftretende Ausgangsspannung. Es erfolgt also keine Verstärkung. Da U_{BE} sehr klein ist, sind die beiden gegenphasigen Aus-

gangsspannungen nur unwesentlich geringer als die Eingangsspannung.

Bei der praktischen Realisierung sind allerdings einige Besonderheiten zu berücksichtigen, vor allem dann, wenn mit dieser Phasenumkehrstufe Gegentakt-B-Verstärker angesteuert werden. Zum Beispiel muß verhindert werden, daß sich bei Aussteuerung die vor den Endtransistoren liegenden Koppelkondensatoren aufladen, deren Ladespannung dann im Rhythmus der Signalfrequenz schwankt und den Arbeitspunkt der Endstufe verschiebt. Außerdem ist zu bedenken, daß die Widerstandsverhältnisse im Kollektor- und Emittterkreis des Transistors der Phasenumkehrstufe sehr unterschiedlich sind. Um diese Einflüsse auszugleichen, sind deshalb noch einige Schaltelemente erforderlich.

Gegentaktverstärker mit Transistoren unterschiedlichen Leitungstyps

In der modernen Schaltungstechnik kommen für Gegentaktverstärker Transistoren unterschiedlichen Leitungstyps zum Einsatz. Diese Schaltungen haben u. a. den Vorteil, daß die Phasenumkehrstufe entfallen kann.

Bevor auf Einzelheiten eingegangen wird, sollen die für die Praxis wichtigsten Grundprinzipien entwickelt werden.

Der im Bild 3.76 dargestellte Leistungsverstärker in Kollektorschaltung läßt über den Widerstand R_E nur einen begrenzten Ausgangsstrom zu. Man erhält wesentlich größere Ausgangsleistung und einen besseren Wirkungsgrad, wenn R_E entsprechend Bild 3.80 durch einen weiteren Emittterfolger ersetzt wird.

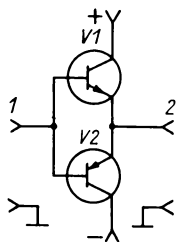


Bild 3.80. Prinzip des Serien-Gegentaktverstärkers mit Transistoren unterschiedlichen Leitungstyps

Beide Endstufentransistoren werden in Kollektorschaltung betrieben.

Bei positiver Eingangsspannung leitet Transistor V_1 , und V_2 sperrt; bei negativer Eingangsspannung ist es umgekehrt. Beide Transistoren sind also abwechselnd je eine Halbperiode leitend. Das ist charakteristisch für Gegentaktbetrieb.

Für höhere Ausgangsströme benötigt man Transistoren mit größerer Stromverstärkung. Solche Transistoren lassen sich aus zwei oder mehr Einzeltransistoren realisieren, wenn sie in Standard-Darlington- oder Komplementär-Darlington-Schaltung zusammengefügt sind (s. Abschn. 3.2.3.).

Werden nun die Transistoren V_1 und V_2 im Bild 3.80 durch Standard-Darlington-Schaltungen entsprechend Bild 3.6 a ersetzt, dann erhält man die im Bild 3.81 wiedergegebene Schaltung.

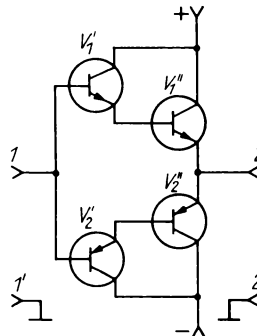


Bild 3.81. Prinzip des Serien-Gegentaktverstärkers mit Standard-Darlington-Schaltungen

Es wurde schon mehrfach auf die Vor- und Nachteile des B-Betriebs hingewiesen. In der Praxis strebt man, deshalb meistens AB-Betrieb an. Man erreicht AB-Betrieb, indem man bei der Schaltung nach Bild 3.80 zwischen die Basisanschlüsse von V_1 und V_2 jeweils eine Gleichspannung legt, damit ein gewünschter Ruhestrom durch V_1 und V_2 fließt. Diese Spannungen kann man durch Spannungsabfälle über in Durchlaßrichtung betriebene Dioden V_3 und V_4 erhalten. Der Wirkungsgrad des Verstärkers sinkt damit unvermeidlich um mehr als 15 % gegenüber B-Betrieb ab.

Ein weiteres Problem ist die Ruhestromstabilisierung. Sie läßt sich durch je einen Widerstand R_1 und R_2 in die Emittterleitung der Endtransistoren erreichen, die eine Stromgegenkopplung bewirken. Da diese jedoch in Reihe mit dem Verbraucher liegen, wird die Ausgangsleistung weiter herabgesetzt. R_1 und R_2 sind daher sehr niederohmig zu dimensionieren. Bei höheren Ansprüchen können sie durch in Durchlaßrichtung betriebene Dioden überbrückt werden.

Bild 3.82 zeigt einen AB-Verstärker mit den besprochenen Schaltungseinzelheiten.

Der Strom durch die Dioden V_3 und V_4 sollte unabhängig von der Aussteuerung konstant bleiben, damit auch der an ihnen auftretende

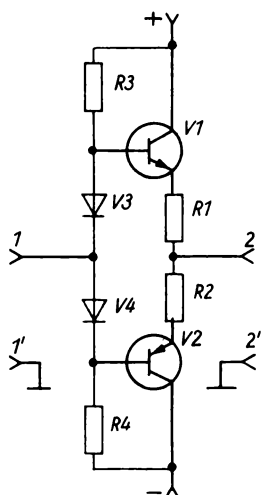


Bild 3.82. Gegentaktverstärker im AB-Betrieb

Spannungsabfall konstant bleibt. Um diese Forderung zu erfüllen, ersetzt man die Widerstände $R3$ und $R4$ durch Konstantstromquellen, die einen von der Eingangsspannung unabhängig konstanten Strom bei sehr hohem Innenwiderstand liefern. Damit muß der Vorverstärker nur den zur Ansteuerung von $V1$ und $V2$ benötigten Basisstrom liefern; der Eingangswiderstand der Schaltung wird also wesentlich erhöht. Bild 3.83 zeigt diese Schaltungsverbesserung.

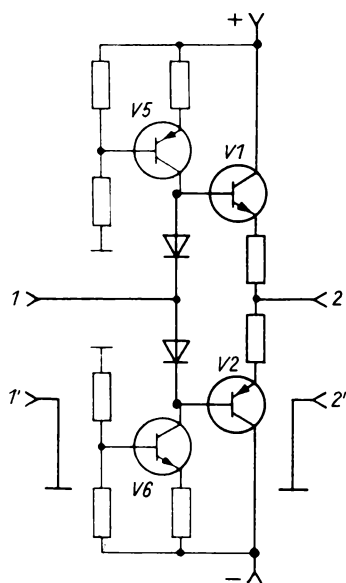


Bild 3.83. AB-Gegentaktverstärker mit Konstantstromquellen

Der im Bild 3.83 entwickelte Endverstärker bringt keine Spannungsverstärkung. Es können sich damit Schwierigkeiten bei der Dimensionierung der Vorstufe ergeben, da diese die gesamte für die Aussteuerung benötigte Spannung liefern muß. Vorteilhaft ist, wenn der Endverstärker selbst eine Spannungsverstärkung $v_u > 1$ aufweist. Realisiert werden kann das dadurch, daß eine der Konstantstromquellen durch einen Transistor in Emitterschaltung ersetzt wird, der damit als Treiberstufe arbeitet. Bild 3.84 zeigt die entsprechende Schaltung.

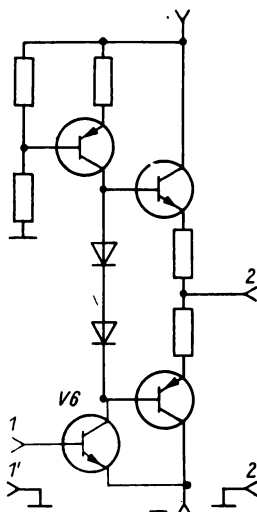


Bild 3.84. AB-Gegentaktverstärker mit Ansteuerschaltung

Die an 1 angelegte Eingangsspannung wird durch $V6$ verstärkt und tritt vergrößert zwischen den beiden Dioden auf. Das Eingangspotential liegt in der Nähe der negativen Betriebsspannung. Hat man nur Wechselspannungen zu verstärken und wünscht man eine besonders große Ausgangsaussteuerbarkeit, dann läßt sich die Konstantstromquelle $V5$ durch einen ohmschen Widerstand ersetzen, der dynamisch durch das Bootstrap-Schaltungsprinzip (s. Abschn. 3.2.2.) vergrößert wird. Man benötigt dazu einen Kondensator $C1$.

Bild 3.85 zeigt die entsprechende Schaltung.

Wirkungsweise der Bootstrap-Schaltung

Steigt beispielsweise durch die Aussteuerung das Kollektorpotential von $V6$, dann steigt auch die Ausgangsspannung um den gleichen Betrag. Diese Spannungsänderung wird durch $C1$ übertragen, so daß auch die Spannung zwischen $R31$ und $R32$ um den gleichen Wert ansteigt. Damit bleibt die Gleichspannung an $R31$ weit-

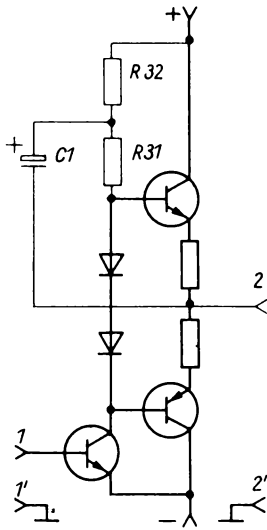


Bild 3.85. AB-Gegentaktverstärker mit Wechselspannungs-Bootstrap

gehend konstant. Diese schon einmal begründete Forderung wird damit erfüllt, und man erreicht so eine große Ausgangssteuerbarkeit für Wechselspannungen. Steht keine Betriebsspannungsquelle mit Mittelanzapfung zur Verfügung, dann hat man die negative Betriebsspannung in der Schaltung nach Bild 3.85 Null zu machen und das Ausgangsruhepotential auf $U_B/2$ einzustellen. Der Verbraucher muß dann über den möglichst großen Koppelkondensator

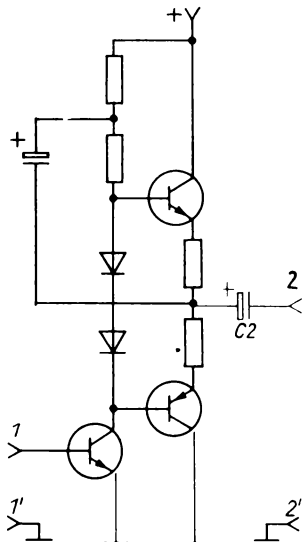


Bild 3.86. Gegentaktverstärker für eine Betriebsspannung

C2 angeschlossen werden. Eine Gleichstromverstärkung ist natürlich nicht mehr möglich. Bild 3.86 zeigt die ausgeführte Schaltung. Endstufen mit höheren Ausgangsströmen werden mit Darlington-Schaltungen realisiert. Ersetzt man die Transistoren $V1$ und $V2$ in der Schaltung nach Bild 3.83 durch die Standard-Darlington-Schaltungen V'_1/V''_1 und V'_2/V''_2 , dann ergibt sich die im Bild 3.87 angegebene Schaltung.

Da in jedem Darlington-Zweig zwei Basis-Emitter-Strecken in Reihe liegen, müssen für die Vorspannungserzeugung auch jeweils zwei Dioden $V31/V32$ bzw. $V41/V42$ eingefügt werden. Zwischen den Emittoren der Transistoren V'_1 und V'_2 ist ein Widerstand $R5$ erforderlich und so zu dimensionieren, daß durch V'_1 und V'_2 ein kleiner Ruhestrom fließt, weil der Basisruhestrom der Endtransistoren V''_1 und V''_2 im allgemeinen zu klein ist. Da die Kollektorströme der Transistoren V'_1 und V'_2 um die Stromverstärkung der Endtransistoren geringer sind, sind auch die Verlustleistungen um den gleichen Faktor geringer. V'_1 und V'_2 sind also im Gegensatz zu V''_1 und V''_2 keine Leistungstransistoren. Alle in Gegentakt geschalteten Transistoren sollten gut in ihren Parametern übereinstimmen.

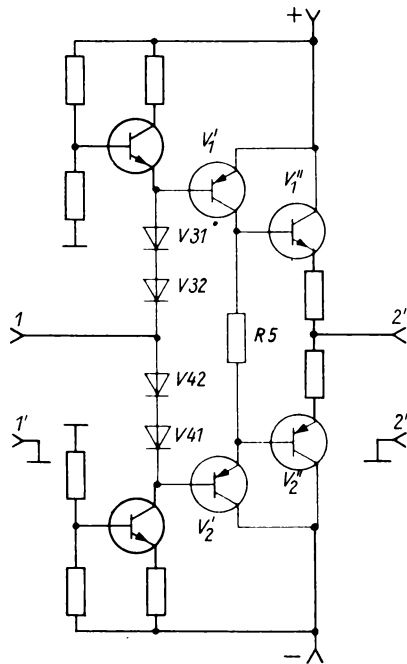


Bild 3.87. Komplementärer Gegentaktverstärker mit Standard-Darlington-Schaltungen

Quasikomplementäre Gegentaktverstärker

Die in Gegentakt geschalteten Transistoren sollen gut in ihren Daten übereinstimmen. Bei Leistungstransistoren unterschiedlichen Leitungstyps ergeben sich dabei technologisch einige Schwierigkeiten. Es ist deshalb oft vorteilhaft, Endtransistoren des gleichen Leitungstyps zu verwenden.

Ersetzt man im Bild 3.80 V_1 durch eine als npn-Transistor wirkende Standard-Darlington-Schaltung und V_2 durch eine als pnp-Transistor wirkende Komplementär-Darlington-Schaltung, dann erhält man das im Bild 3.88 wiedergegebene Prinzip des quasikomplementären Gegentaktverstärkers. Der Vorteil dieses Schaltungsprinzips besteht darin, daß die Endtransistoren den gleichen Leitungstyp haben.

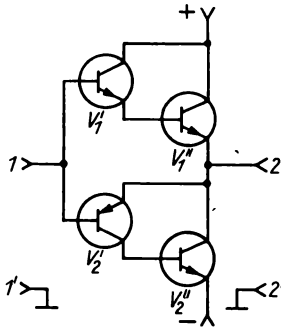


Bild 3.88. Prinzip des quasikomplementären Gegentaktverstärkers

Ersetzt man die untere Standard-Darlington-Schaltung V_1'/V_1'' in der Schaltung nach Bild 3.87 durch eine Komplementär-Darlington-Schaltung, dann ergibt sich die im Bild 3.89 wiedergegebene quasikomplementäre Gegentaktendstufe.

Die Vorspannung der Komplementär-Darlington-Schaltung braucht dann nur mit einer Diode V_4 bereitgestellt zu werden, da nur eine Basis-Emitter-Strecke vorliegt. Die Widerstände R_51 und R_52 ermöglichen einen definierten Ruhestrom durch V_1' und V_1'' . Es gibt nun viele Schaltungsmodifikationen, um speziellen Anforderungen gerecht zu werden. Sie lassen sich praktisch alle auf die dargestellten Grundprinzipien zurückführen.

Schutzschaltungen in Gegentaktverstärkern

Bei zusätzlichen Schaltungsmaßnahmen ist besonders häufig die elektronische Ausgangsstrombegrenzung anzutreffen. Sie dient zum Schutz der teuren Leistungstransistoren bei Überlastung bzw. Kurzschluß. Eine mögliche

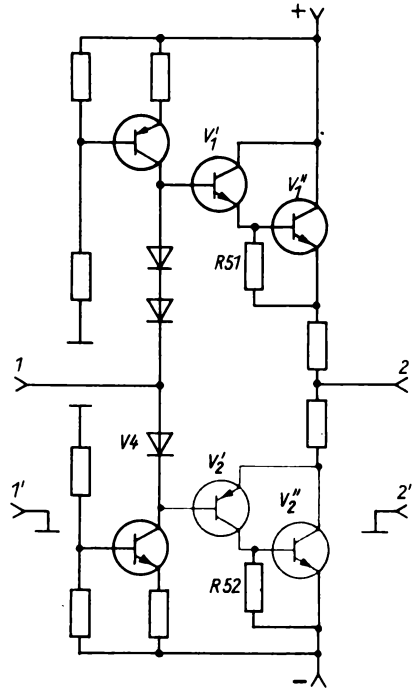


Bild 3.89. Quasikomplementärer Gegentaktverstärker

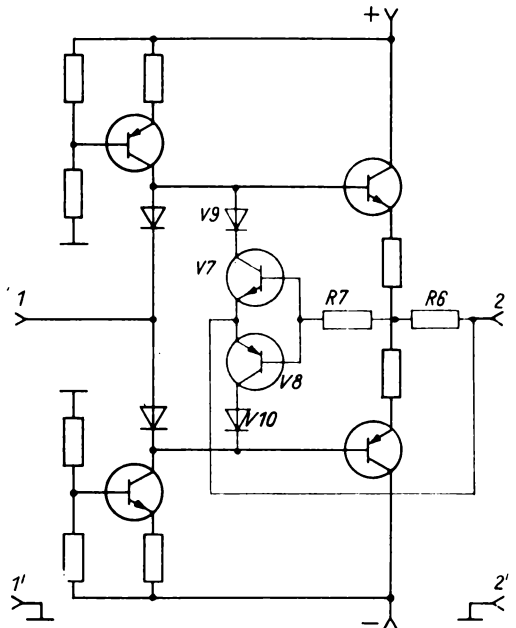


Bild 3.90. Elektronische Ausgangsstrombegrenzung im Gegentaktverstärker

Schaltungsvariante wird dazu im Bild 3.90 dargestellt.

Die über R_6 abfallende Spannung ist dem Ausgangsstrom proportional. Überschreitet dieser einen zulässigen Höchstwert, werden die Strombegrenzungstransistoren V_7 und V_8 leitend, und die Basisströme von V_1 und V_2 werden kleiner. Der Ausgangsstrom kann damit nicht weiter ansteigen. Die Dioden V_9 und V_{10} sind notwendig, damit bei gesperrtem Endtransistor die Kollektor-Basis-Dioden von V_7 und V_8 nicht leitend werden können. Der Widerstand R_7 soll die Strombegrenzungstransistoren vor zu hohen Basisstromspitzen schützen.

3.6.4.

Endverstärker mit integrierten Schaltkreisen

Endverstärker werden auch als integrierte Schaltkreise ausgeführt.

Die Innenschaltungen dieser Schaltkreise sind zwar unterschiedlich, haben aber alle stark vereinfacht folgende Grundstrukturen:

- **Eingangsstufe**
oft als Differenzverstärkerstufe ausgeführt, mit Konstantstromquelle als Arbeitswiderstand der ersten Stufe
- **Treiberstufe**
zur Aussteuerung der Endstufe
- **Gegentaktendstufe**,
die entweder als komplementärer oder quasi-komplementärer Gegentaktverstärker ausgeführt ist
- **Stromversorgungsteil und Regelungsnetzwerk**, letzteres zur Stabilisierung der Mittenspannung auf etwa die halbe Betriebsspannung
- **Schutzschaltungen**
zur Temperatur-, Strom- und Leistungsbegrenzung sowie gegen Überspannungen bei Endverstärkern mit Ausgangsleistungen $> 5 \text{ W}$

Typische integrierte Endverstärker sind beispielsweise die vom VEB Halbleiterwerk Frankfurt (Oder) gefertigten 1 W-NF-Verstärker A 211, 6 W-NF-Verstärker A 210, 16 W-NF-Verstärker A 2030 und die $2 \times 5 \text{ W}$ sowie $2 \times 10 \text{ W}$ Doppel-NF-Verstärker A 2000 und A 2005.

Die Innenschaltung der Schaltkreise interessiert den Anwender erst in zweiter Linie; für ihn sind vor allem die von den Herstellern angegebenen Grenzwerte, elektrischen Kennwerte und allgemeinen Applikationshinweise wichtig.

Mit diesen Angaben ist man in der Lage, den praktischen Bedürfnissen entsprechende Schaltungen zu realisieren.

Als Beispiel sei im Bild 3.91 ein mit dem Schaltkreis A 211 konzipierter 1 W-NF-Verstärker wiedergegeben. Die vom Hersteller angegebenen Anschlußbelegungen sind

1	Bootstrap	8	Eingang
2	Betriebsspannung	9	Gegenkopplung
3,4,5	Masse	10,11,12	Masse
6	Ausgang	13,14	Frequenzkompensation
7	Masse		

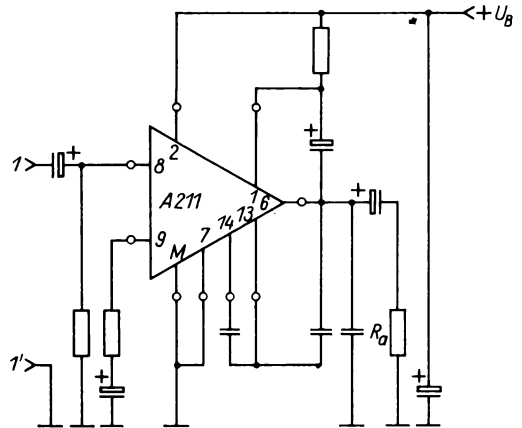


Bild 3.91. NF-Verstärker mit dem integrierten Schaltkreis A 211

Maximale Betriebsspannung, maximale Eingangsspannung, Standardbeschaltung der Frequenzkompensation u. a. sind ebenfalls den Herstellerangaben zu entnehmen. Im Interesse eines hohen Ausgangsspannungshubs und kleinen Klirrranteils des oberen Teiles des Ausgangssignals wird Wechsellspannungs-Bootstrap durch Belegung des Anschlusses 1 angewandt. Müssen bei größeren Ausgangsleistungen andere spezielle Parameter erreicht werden, dann läßt sich wie folgt verfahren. Es wird eine Leistungsstufe mit diskreten Transistoren aufgebaut, die von einem integrierten Schaltkreis geringerer Ausgangsleistung angesteuert wird. Bild 3.92 zeigt eine derartige Schaltung mit dem NF-Verstärker A 211, der als Vorverstärker und Treiber arbeitet und die komplementäre Endstufe aussteuert.

Weil dabei mehr Schaltelemente erforderlich sind, ist der Aufwand größer und die Zuverlässigkeit

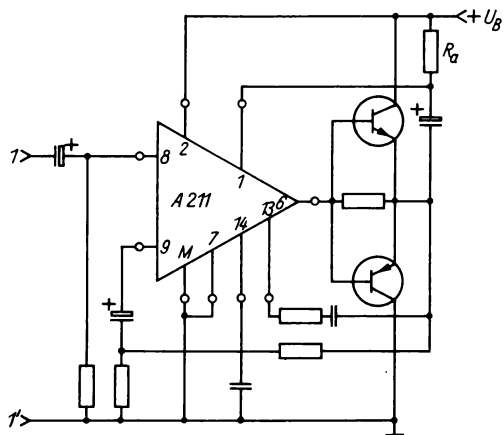


Bild 3.92. Integrierter NF-Verstärker mit aus diskreten Elementen aufgebauter Endstufe höherer Leistung

sigkeit kleiner als beim Einsatz eines geeigneten IS. Derartige Schaltungen sind also nur dann gerechtfertigt, wenn für spezielle Anforderungen keine entsprechenden integrierten Schaltkreise zur Verfügung stehen.

Die Fortschritte in der Halbleitertechnik haben die Entwicklung integrierter Endverstärker ermöglicht, die mehr als 25 W Ausgangsleistung liefern können. Diese Schaltkreise werden vor

allem in Geräten der Unterhaltungselektronik (Rundfunk-, Fernseh- und Fonogeräten) eingesetzt, sind aber auch für Motorsteuerungen, als Leistungsspannungswandler u. a. geeignet.

Den typischen Aufbau derartiger IS zeigt der Übersichtsplan des integrierten Doppel-NF-Endverstärkers A 2000 (A 2005) im Bild 3.93.

Der Schaltkreis hat einen invertierenden und einen nichtinvertierenden Eingang. Der Vorverstärker wird durch die integrierte Kapazität C im oberen Frequenzbereich gegengekoppelt; es kann dadurch eine äußere Frequenzkompensation entfallen. Die Treiberstufe hat einen Bootstrapschluß. Damit kann durch externe Kondensatoren die Spannung der Treiberstufe aufgestockt und eine niedrige Sättigungsspannung von unter 2 V der Endstufentransistoren erreicht werden. Die Gegentakt-B-Endstufen können Ausgangsspitzenströme von 2,5 A (A 2000) bzw. 3,5 A (A 2005, A 2030) liefern. Bedingt durch den geregelten Endstufenruhestrom treten auch im unteren Aussteuerbereich keine Übernahmeverzerrungen auf.

Die Gehäuse dieser IS haben eine sehr hohe Wärmeleitfähigkeit. Bei sachgerechter Montage brauchen die externen Kühlkörper nur für die üblichen Betriebstemperaturen dimensioniert zu werden.

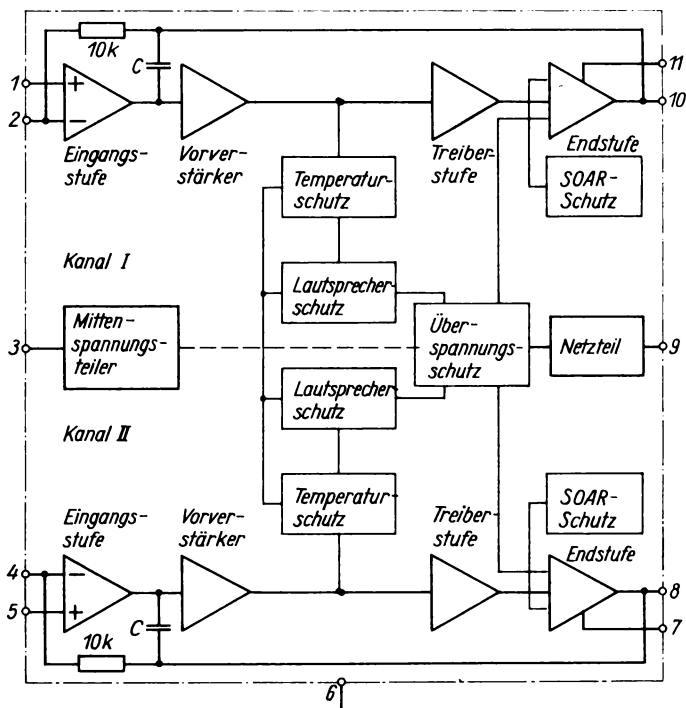


Bild 3.93. Integrierter Doppel-NF-Endverstärker; Übersichtsplan

Die zahlreichen *Schutzschaltungen* dienen ausschließlich dem Schutz des Schaltkreises, nicht etwa dem anderer Schaltungsteile. Sie sorgen für eine hohe Betriebszuverlässigkeit und sollen etwas näher betrachtet werden:

- **Temperaturschutz.** Ein als Temperatursensor arbeitender Transistor ist auf dem Chip direkt am Ausgangstransistor positioniert. Erhöht sich die Temperatur, dann ändert sich die Basis-Emitterspannung des Sensors, um bei etwa 150 °C die Ansteuerung der Endstufe zu vermindern.
- **Überspannungsschutz.** Spannungsspitzen auf den Betriebsspannungsleitungen in Kraftfahrzeugen gefährden die dort eingesetzten Bauelemente. Durch die Schutzschaltung wird bei kurzzeitigen Überspannungen die Treiberstufe abgeschaltet und die Endstufe geschützt.
- **Lautsprecherkurzschlußschutz.** Ein Überstrom bzw. Kurzschluß der Ausgänge gegen Masse verursacht an den betreffenden Bonddrähten zwischen dem Chip und den Gehäuseanschlüssen einen dem Überstrom proportionalen Spannungsabfall. Dieser wird in der Schutzschaltung verstärkt und die Ansteuerung des betr. Endstufenzweiges reduziert. Wenn der Kurzschluß beseitigt ist, ist die Schaltung wieder betriebsbereit.
- **SOAR-Schutz.** SOAR (safe operating area) bedeutet sicherer Arbeitsbereich (für Leistungstransistoren). Diese Schutzschaltung reduziert die Ansteuerung der Endstufen, wenn deren Aussteuerung oberhalb der Verlustleistungshyperbel erfolgt. Es wird dabei ständig die Spannung über den Ausgangstransistoren kontrolliert; je höher diese Spannung ist, desto niedriger muß der Kollektorstrom werden.

Die Funktionsgruppe *Netzteil* besitzt eine Bandgap-Referenzspannungsquelle (s. Abschn. 3.2.3). Dadurch läßt sich der Schaltkreis noch bis zu einer sehr kleinen Betriebsspannung von etwa 4 V betreiben.

Der *Mittenspannungsteiler* ist so bemessen, daß die Ausgangsmittenspannung symmetrisches Aussteuern zuläßt, wenn keine Bootstrapkondensatoren angeschaltet sind. Bei Bootstrap-Betrieb muß die Ausgangsmittenspannung durch R_4 geringfügig angehoben werden.

Wird der herausgeführte Anschluß des Mittenspannungsteilers auf Masse gelegt, dann folgt dem auch die Ausgangsmittenspannung. Der

Schaltkreis wird dadurch stummgeschaltet bei gleichzeitig starker Reduzierung der Ruhestromaufnahme.

Die externe Beschaltung dieser integrierten Schaltkreise ist problemlos; es sind aber die von den Herstellern angegebenen Applikationshinweise zu beachten. Prinzipiell gilt, daß die Betriebsspannung direkt an den Schaltkreisanschlüssen abzublocken und in jedem Fall parallel zum Lastwiderstand ein Boucherot-Glied zu schalten ist. Das Boucherot-Glied ist eine Reihenschaltung aus Kondensator mit $C = 100 - 220 \text{ nF}$ und Widerstand mit $R = 1 \Omega$. Durch dieses Glied werden Eigenschwingungen der Schaltung bei hohen Frequenzen verhindert.

Bild 3.94 zeigt eine typische Grundsaltung, die mit dem IS A 2030 realisierbar ist. Die vom Hersteller angegebenen Anschlußbelegungen sind

- 1 nichtinvertierender Eingang
- 2 invertierender Eingang
- 3 negative Betriebsspannung
- 4 Ausgang
- 5 positive Betriebsspannung.

Da im Beispiel mit nur einer Betriebsspannung + U_B gearbeitet werden soll, muß der nichtinvertierende Eingang über den Spannungsteiler $R_3 - R_4$ und den Widerstand R_5 auf $U_B/2$ gelegt werden. Die Anschaltung des Lastwiderstandes R_a (Lautsprecher) erfolgt über den Kopplerkondensator C_7 , dessen Kapazität bei 2200 μF liegt. Die untere Grenzfrequenz wird von den Größen C_7 und R_a bestimmt; bei größerem R_a kann deshalb C_7 bis auf 1000 μF verringert werden.

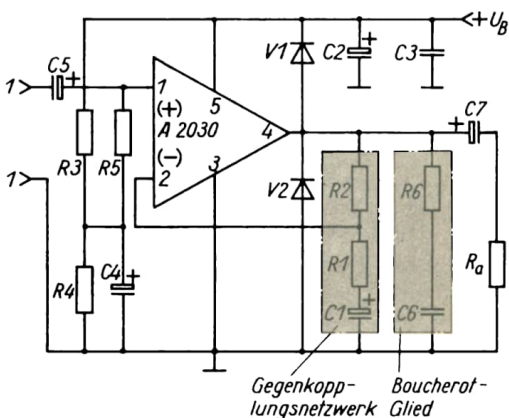


Bild 3.94. NF-Endverstärker mit dem integrierten Schaltkreis A 2030

gert werden. Das Boucherot-Glied $R6/C6$ verhindert das Eigenschwingen der Schaltung. Mit dem Gegenkopplungsnetzwerk $R1, R2$ und $C1$ wird die gewünschte Verstärkung eingestellt; sie ergibt sich aus $v'_u = 1 + \frac{R_2}{R_1}$. $V1$ und $V2$ sind sehr schnelle Dioden (z. B. SAY 12) und schützen den Schaltkreisausgang vor induktiven Spannungsspitzen. Mit $C2$ und $C3$ werden unerwünschte Rückwirkungen über die Betriebsspannung abgeblockt. $C4$ ist ein Siebkondensator für den Mittenspannungsteiler. $C5$ dient zur Gleichspannungstrennung zwischen dem Vorstufenausgang und dem Eingang des A 2030. Die erreichbare Ausgangsleistung für sinusförmige Signale läßt sich überschlagsmäßig durch $P_2 = 0,5 U_B^2/R_a$ bestimmen. Im Interesse eines niedrigen Klirrfaktors wird die Leistung aber stets unter 70 % des mit obiger Gleichung ermittelten Wertes P_2 liegen.

Beispiel

Welche Ausgangsleistung kann ein Endverstärker liefern, der mit einer Spannung von 14 V betrieben wird und auf einen Lastwiderstand mit 4 Ω arbeitet?

Gegeben:

$$U_B = 14 \text{ V}$$

$$R_a = 4 \Omega$$

Gesucht:

$$P_2$$

Lösung:

$$P_2 \leq 0,7 P'_2$$

$$P'_2 = 0,5 U_B^2/R_a = 0,5 \cdot 14^2 \text{ V}^2/4 \Omega = 24,5 \text{ W}$$

$$P_2 \leq 17 \text{ W}$$

Die maximal erreichbare Ausgangsleistung bei sinusförmiger Aussteuerung liegt bei 17 W.

Um höhere Ausgangsleistungen bei niedriger Betriebsspannung zu erreichen, wendet man die Brückenschaltung mit zwei gleichen Verstärkern an. Unter gleichen Betriebsbedingungen erreicht man mit *Brückenverstärkern* gegenüber Einzelverstärkern etwa die vierfache Ausgangsleistung. Zu beachten ist aber, daß der Maximalwert des vom Lastwiderstand begrenzten Ausgangsstroms nicht überschritten wird. Der untere Grenzwert des Lastwiderstandes ist

$$R_{\text{amin}} = \frac{U_B - 2 U_{\text{CESat}}}{I_{2\text{max}}}$$

Bei $U_B = \pm 9 \text{ V}$, $U_{\text{CESat}} \approx 2 \text{ V}$ bei $I_{2\text{max}} = 3,5 \text{ A}$ muß also der Lastwiderstand R_a mindestens 4 Ω betragen.

Da die Einzelverstärker jeweils einen invertierenden und einen nichtinvertierenden Eingang haben, sind für den Brückenverstärker keine zu-

sätzlichen Phasenumkehrstufen erforderlich, und der Aufbau wird sehr einfach.

Bild 3.95 zeigt die Schaltung eines Brückenverstärkers mit zwei IS A 2030. Die gegenphasige Ansteuerung erfolgt so, daß der rechte Verstärker sein Eingangssignal auf den invertierenden Eingang vom Ausgang des linken Verstärkers über den Widerstand $R7$ erhält.

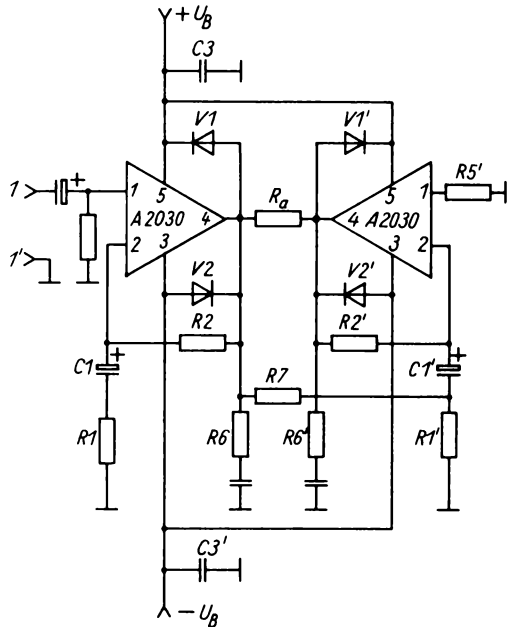


Bild 3.95. Brückenverstärker mit zwei IS des Typs A 2030

Die im Bild 3.95 wiedergegebene Schaltung arbeitet mit positiver und negativer Betriebsspannung. Die fehlenden Auskoppelkondensatoren an den Ausgängen entfallen aber auch dann, wenn nur mit einer Betriebsspannung gearbeitet wird. Das erfordert aber, daß die Ausgangsmittenspannung der beiden Schaltkreise übereinstimmen, weil sonst ein Gleichstrom durch den als Lastwiderstand wirkenden Lautsprecher fließen würde. Die Differenz der Ausgangsmittenspannung darf höchstens 100 mV betragen; diese geringe Spannung darf an einen 4 Ω -Lautsprecher ohne dessen Gefährdung dauerhaft anliegen.

Die Schaltungsdetails im Bild 3.94 treten hier auch wieder auf und haben die gleichen Aufgaben. Das sind die Boucherot-Glieder, Gegenkopplungsnetzwerke, Schutzdioden, Kondensatoren an den Betriebsspannungsanschlüssen und ein Eingangskondensator zur Gleichspannungstrennung der Vorstufe.

Bild 3.96 zeigt die Schaltung eines Stereo-Endverstärkers. Als Schaltkreis wird der Typ A 2000 oder A 2004 eingesetzt, dessen Übersichtsplan im Bild 3.93 angegeben ist. Beide Kanäle sind völlig gleich aufgebaut; die zum rechten Kanal gehörenden Bauelemente werden durch einen hochgestellten Strich gekennzeichnet.

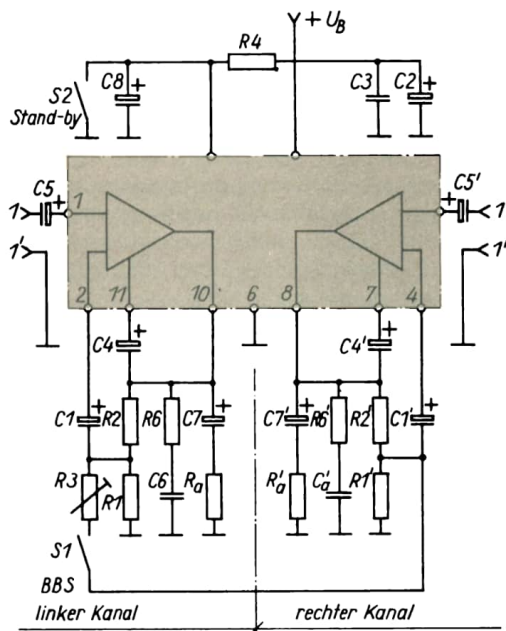


Bild 3.96. Stereo-Endverstärker mit dem Schaltkreis A 2000

Die schon weiter vorn beschriebenen Schaltdetails treten hier auch wieder auf. Mit dem Gegenkopplungsnetzwerk R_1 , R_2 und C_1 wird die Spannungsverstärkung auf das gewünschte Maß gesenkt. Sie ergibt sich aus $v_u = 1 + R_2/R_1$ und läßt sich zwischen 24 – 52 dB einstellen. Um ausreichende Stabilität zu erreichen, muß der Spannungsteiler R_1 , R_2 niederohmig sein. Die Anschaltung des Lastwiderstandes R_L (Lautsprecher) erfolgt über den Auskoppelkondensator C_7 . Parallel zum Ausgang liegt das Boucherot-Glied $R_6 - C_6$. Durch den Bootstrap-Kondensator C_4 erfolgt eine Spannungsaufstockung; bei Bootstrapbetrieb muß mit R_4 die Ausgangsmittenspannung geringfügig angehoben werden. Mit C_2 und C_3 werden unerwünschte Rückwirkungen über die Betriebsspannung abgeblockt. Durch C_8 werden Störspannungen unterdrückt. Der Kondensator

C_5 dient zur Gleichspannungstrennung zwischen dem Vorstufenausgang und dem Eingang des A 2000.

Über den Einstellwiderstand R_3 läßt sich die Stereobasisbreite verringern oder verbreitern. Die Basisbreite ist der Lautsprecherabstand einer Stereoanlage und sollte in Wohnräumen etwa 3 m betragen. Ist dieser Abstand geringer, so wird auch der Stereo-Höreindruck zusammengedrängt. Um diesen Mangel auszugleichen, wird jedem Kanal ein Anteil des jeweils anderen, der in der Phase um 180° gedreht wurde, zugefügt. Damit verschiebt sich die Basisbreite scheinbar über die Lautsprecher hinaus. Mit dem Schalter S_1 wird die Basisbreitensteuerung (BBS) wahlweise zu- oder abgeschaltet.

Mit dem Schalter S_2 läßt sich die Mittenspannung gegen Masse kurzschließen; damit wird der Schaltkreis stummgeschaltet und es ist eine Stand-by-Beschaltung möglich. Allgemein bezeichnet man mit Stand-by eine Anzeigeeinrichtung, die die Betriebsbereitschaft elektronischer Geräte und Einrichtungen signalisiert.

Die Schaltung arbeitet mit nur einer Betriebsspannung im Bereich von +4 bis +18V.

Bei entsprechender Dimensionierung wird eine Bandbreite von 20 Hz bis > 20 kHz erreicht. Die für Stereoverstärker gestellten Forderungen an die Kanaltrennung werden problemlos erfüllt.

3.6.5.

Endverstärker mit sehr hohem Wirkungsgrad (D-Verstärker)

Endverstärker für sehr hohe Ausgangsleistungen sollen einen wesentlich größeren Wirkungsgrad haben als den mit Gegentakt-B-Schaltungen erreichbaren. Dafür ist vor der Verstärkung eine Analog-Digital-Signalwandlung mit abschließender Rückwandlung erforderlich.

Das zu verstärkende analoge Signal wird nach Vorverstärkung in einem Modulator in eine entsprechende dauermodulierte Impulsfolge umgewandelt (s. Abschn. 6.1.2.). Im anschließenden Verstärker können diese Impulse mit sehr hohem Wirkungsgrad verstärkt werden, da die Transistoren im Schalterbetrieb arbeiten. Die abschließende Rückwandlung in ein analoges Signal erfolgt durch einen einfachen Tiefpaß, indem der Mittelwert der Impulsfolge gebildet wird.

Bild 3.97 zeigt den Übersichtsschaltplan des D-Verstärkers.

Der Vorverstärker wird als Differenzverstärker aufgebaut. Pulsdauermodulator und Impuls-generator werden zusammengefaßt durch einen astabilen Multivibrator (s. Abschn. 4.3.2.) realisiert, dessen Impulsdauer durch den Vorverstärker gesteuert wird. Als Impulsfrequenz wählt man das 5- bis 7fache der höchsten Signalfrequenz, bei Tonfrequenzverstärkern meist 100 kHz.

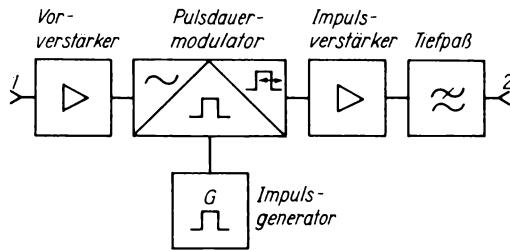


Bild 3.97. Übersichtsschaltplan eines D-Verstärkers

Als Impulsverstärker eignet sich eine komplementäre Gegentaktschaltung; ihre obere Grenzfrequenz muß 200- bis 400mal höher als die höchste Signalfrequenz sein.

Im abschließenden Tiefpaß werden alle Schwingungen, die oberhalb der höchsten Signalfrequenz liegen, gesperrt und der Mittelwert der Impulse gebildet. Am Ausgang steht dann wieder ein analoges, aber mit sehr hohem Wirkungsgrad verstärktes Signal zur Verfügung. Gegenüber dem Gegentakt-B-Verstärker hat der D-Verstärker folgende Vor- und Nachteile:

Vorteile

- größerer Wirkungsgrad
- einfachere Arbeitspunkteinstellung und -stabilisierung
- geringere Anforderungen an Bauelementetoleranzen und Linearität der Endstufe.

Nachteile

- größerer Schaltungsaufwand
- große Bandbreite des Impulsverstärkers
- Gefahr einer Störstrahlung wegen der energiereichen steilen Impulse
- im Modulator sind Kondensatoren, im Tiefpaß Spule und Kondensator erforderlich. Diese Blindschaltelemente können in der integrierten Technik nur mit Schwierigkeiten realisiert werden.

Wegen dieser Nachteile kann der D-Verstärker nicht universell als Endverstärker eingesetzt werden.

Zusammenfassung

Endverstärker sind Verstärker, bei denen eine hohe Ausgangsleistung angestrebt wird, während die Spannungsverstärkung eine zweitrangige Rolle spielt. Man unterscheidet je nach Lage des Arbeitspunktes A-, AB-, B- und C-Betrieb.

Eintaktverstärker arbeiten stets im A-Betrieb. Der Arbeitspunkt liegt dabei etwa auf der Mitte des „geradlinigen“ Teiles der Steuerkennlinie. Zur Dimensionierung der Schaltung legt man das Ausgangskennlinienfeld zugrunde und trägt die Aussteuerungsgrenzen ein; diese sind gegeben durch die maximale Kollektorverlustleistung, die Grenzwerte für Strom und Spannung sowie Restspannung und Kollektorreststrom. Gegentaktverstärker können in allen Betriebsarten arbeiten. Gegenüber Eintaktverstärkern liefern sie bei kleinen nichtlinearen Verzerrungen größere Ausgangsleistungen.

Gegentaktverstärker können als Parallel- und Serien-Gegentaktschaltung mit Transistoren gleichen oder unterschiedlichen Leitungstyps aufgebaut werden. In der modernen Schaltungstechnik dominieren Serien-Gegentaktverstärker. Bei diesen sind die beiden Endtransistoren gleichstrommäßig in Reihe, dagegen wechselstrommäßig parallelgeschaltet. Bevorzugt werden Transistoren unterschiedlichen Leitungstyps eingesetzt. Als Grundstrukturen ergeben sich damit komplementäre Gegentaktverstärker, komplementäre Gegentaktverstärker mit Standard-Darlington-Schaltungen und quasikomplementäre Gegentaktverstärker. Derartige Verstärker brauchen nur gleichphasig angesteuert zu werden, so daß eine Phasenumkehrstufe überflüssig ist.

Endverstärker werden auch als integrierte Schaltkreise ausgeführt. Für höhere Ausgangsleistungen werden z. Z. häufig aus diskreten Elementen aufgebaute Endstufen realisiert, die durch einen als Vorverstärker und Treiberstufe arbeitenden integrierten Schaltkreis angesteuert werden.

Endverstärker mit sehr hohem Wirkungsgrad sind D-Verstärker. Bei diesen wird nach einer Vorverstärkung das analoge Signal in eine dauermodulierte Impulsfolge umgewandelt, verstärkt und in einem abschließenden Tiefpaß wieder in ein analoges Signal rückgewandelt.

Neben großer Vorteile haben D-Verstärker einige nicht zu übersehende Nachteile, so daß einem universellen Einsatz Grenzen gesetzt sind.

3.7.

Einstellglieder im Niederfrequenzverstärker

3.7.1.

Verstärkungseinsteller

Die Verstärkung eines NF-Verstärkers ist bei gegebener Größe des Eingangssignals bestimmend für die Lautstärke der Sprach- und Musikwiedergabe.

Im einfachsten Fall läßt sich die Lautstärke dadurch verändern, daß entsprechend Bild 3.98 in den Verstärkereingang ein verstellbarer Widerstand (Potentiometer) eingeschaltet wird. Die Gesamtverstärkung wird dabei stets verringert.

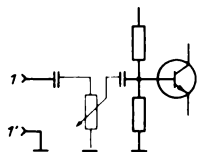


Bild 3.98: Verstärkungseinstellung

Das menschliche Ohr hat eine nichtlineare Empfindlichkeitskurve. Die von ihm empfundene Lautstärkeänderung macht bei großen Lautstärken eine wesentlich größere Änderung der Schalleistung erforderlich als bei kleinen. Aus diesem Grund muß das zur Verwendung kommende Potentiometer eine entsprechende Widerstandsänderung aufweisen, wobei der Widerstand in Abhängigkeit vom Drehwinkel zuerst langsam und dann schnell zunimmt, d. h. einen logarithmischen Verlauf hat.

Die Empfindlichkeitskurve des Ohres ist außerdem frequenzabhängig, was sich darin ausdrückt, daß Schwingungen mit tiefen und hohen Frequenzen bei kleiner Lautstärke mehr als die mit mittleren zurücktreten. Man verwendet deshalb oft eine frequenzabhängige oder gehörrichtige Lautstärkeneinstellung nach Bild 3.99. Das Potentiometer ist mit einem Abgriff versehen, der an eine RC-Reihenschaltung führt. Kommt der Schleifer in die Nähe des Abgriffs oder darunter (kleinere Lautstärke), werden durch das RC-Glied die Eingangsgrößen mit

mittleren Frequenzen stärker geschwächt als jene mit tiefen.

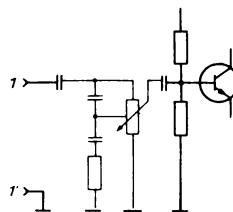


Bild 3.99: Gehörrichtige oder physiologische Lautstärke-einstellung

3.7.2.

Klangeinsteller

Oft ist es wünschenswert, das Klangbild in gewissen Grenzen verändern zu können, indem man beispielsweise Signale mit tiefen oder (und) hohen Frequenzen gegenüber den mit mittleren bevorzugt wiedergibt. Aus der Vielzahl der entwickelten Schaltungen wird im Bild 3.100 ein bewährtes Klangeinstellnetzwerk mit getrennten Höhen- und Tiefeneinstellern angegeben, das zwischen zwei Verstärkerstufen gelegt wird.

Anhand der darübergezeichneten Teilschaltungen wird die Wirkungsweise des Netzwerks verständlich. Mit dem Potentiometer R_2 kann man die Signale mit hohen, mit dem R_4 die mit tiefen Frequenzen einstellen.

Steht beispielsweise der Schleifer des Höheinstellers R_2 oben, ergibt sich der nach Bild 3.100 (I) gezeichnete Hochpaß, der höherfrequente Eingangssignale bevorzugt zur nächsten Verstärkerstufe überträgt.

Steht der Schleifer von R_2 unten, dann erhält man einen Tiefpaß nach Bild 3.100 (II), der die Schwingungen mit hohen Frequenzen kurzschließt.

Der Tiefeneinsteller R_4 liefert bei obenstehendem Schleifer zusammen mit den anderen Schaltelementen einen Tiefpaß nach Bild 3.100 (III) und bewirkt damit eine Baßanhebung. Bei untenstehendem Schleifer von R_4 erhält man einen Hochpaß entsprechend Bild 3.100 (IV), der die Signale mit tiefen Frequenzen benachteiligt.

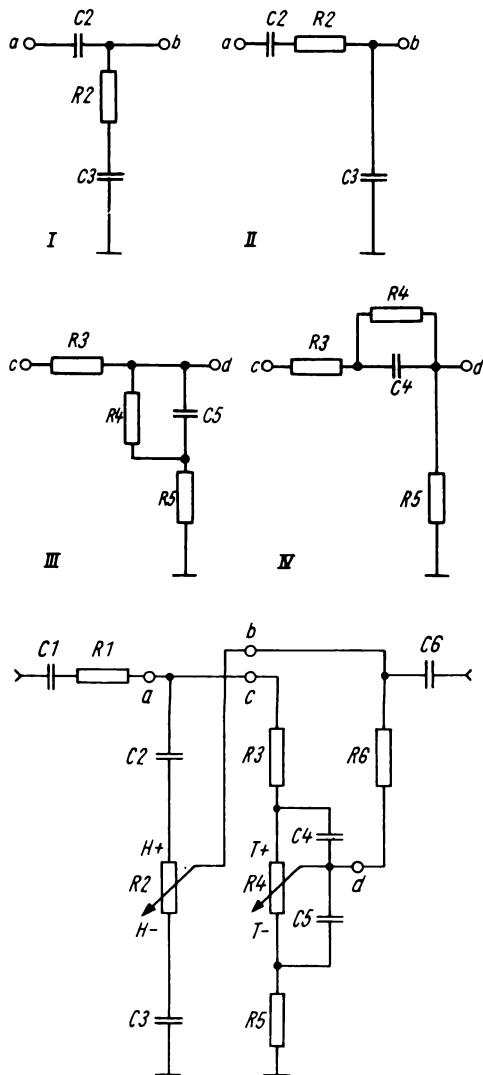


Bild 3.100. Klangeinstellnetzwerk

3.7.3.

Elektronische Einstellglieder

In vielen Fällen sollen die notwendigen Einstellungen in Verstärkern elektronisch mittels einer variablen Gleichspannung erfolgen. Das bringt nicht nur einen höheren Bedienungskomfort der Geräte, sondern hat vor allem die folgend genannten konstruktiven Vorteile:

- Statt mechanisch gekoppelter Potentiometer können Einfachpotentiometer verwendet werden.

- Die Leitungen brauchen nicht abgeschirmt zu werden, da sie nur eine unkritische Gleichspannung übertragen. Damit können die Einstellorgane an der bedienungsgünstigsten Stelle angeordnet werden.

Eine Fernbedienung, drahtlos oder über abgeschirmte Leitungen, ist somit problemlos realisierbar.

Der VEB Halbleiterwerk Frankfurt (Oder) fertigt für elektronisch einstellbare Verstärker zwei integrierte Schaltkreise unter der Typenbezeichnung A 273 und A 274. Diese sind für NF-Stereoverstärker konzipiert, lassen sich aber selbstverständlich auch für Monoverstärker verwenden.

Die Schaltkreise müssen im Verstärker an der elektrisch günstigsten Stelle angeordnet werden, um den Vorteil nichtabgeschirmter Leitungen nutzen zu können.

Der A 273 dient zur Verstärkungs-(Lautstärke-)Einstellung. Dafür sind am integrierten Schaltkreis zwei Möglichkeiten vorgesehen. Ein Stellingang wirkt als Lautstärkeeinsteller gleichsinnig auf die beiden Stereokanäle, entspricht also der Wirkung eines Tandempotentiometers. Der zweite Stellingang bewirkt eine gegensinnige Verstärkungsänderung der beiden Kanäle, entspricht also der Wirkung des Stereobalanceeinstellers. Außerdem enthält der A 273 eine wahlweise abschaltbare gehörrichtige (physiologische) Lautstärkeeinstellung.

Der A 274 dient zur Klangeinstellung in Stereoverstärkern. Der Schaltkreis enthält elektronische Potentiometer und Verstärker für je einen Tiefen- und Höheneinsteller in jedem Kanal. Der Frequenzgang wird dabei durch eine veränderbare frequenzabhängige Gegenkopplung beeinflusst.

Bei elektronischen Einstellgliedern läßt sich die Stromverteilungssteuerung im Differenzverstärker vorteilhaft nutzen. Wie schon im Abschn. 3.2.3. erläutert, wird der von einer Konstantstromquelle gelieferte Gesamtstrom je nach Größe der Eingangsspannungen auf die beiden Zweige des Differenzverstärkers aufgeteilt. Bild 3.101 zeigt das Prinzip.

Die zu verstärkende Spannung wird an den Eingang 13 gelegt. Damit überlagert man den von der Stromquelle mit $V3$ gelieferten Strom I_3 mit dem Eingangssignal. Die zwischen 11 und 12 liegende Steuergleichspannung U_{st} verteilt den Strom I_3 auf den linken und rechten Zweig des Differenzverstärkers. Über $V1$ wird die Ausgangsspannung ausgekoppelt. Da je nach Größe

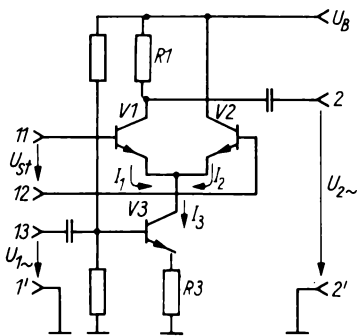


Bild 3.101. Prinzip der Stromverteilungssteuerung im Differenzverstärker

der Steuergleichspannung der Strom I_1 zwischen Null und I_3 liegen kann, läßt sich die Verstärkung $v = U_{2\sim}/U_{1\sim}$ im Bereich $v_{\min} = 0$ bis $v_{\max} = R_1/R_3$ verändern.

Um auftretende Verzerrungen klein zu halten, verwendet man in der Praxis keinen einfachen, sondern einen kreuzgekoppelten Differenzverstärker, auf dem die Stromquelle und ein Eingangstransistor wirken (vgl. Bild 6.21).

Für die beiden funktionell aufeinander abgestimmten Schaltkreise A 273 und A 274 hat der Hersteller ausführliche Applikationshinweise veröffentlicht. Diese sind genau zu beachten, um optimale Ergebnisse zu erreichen. Trotz der notwendigen umfangreichen Außenbeschaltung dieser Schaltkreise mit diskreten Bauelementen sind die eingangs erwähnten Vorteile so groß, daß bei hochwertigen Verstärkeranlagen auf elektronische Einstellglieder nur noch selten verzichtet wird.

Für die Einstellungen der Verstärkung (Lautstärke), der Balance, der Höhen und der Tiefen beim A 273 bzw. A 274 sind 10-k Ω -Einfachpotentiometer mit linearer Stellcharakteristik vorgesehen.

Außerdem sind mehrere, z. T. relativ umfangreiche RC-Netzwerke erforderlich. Mit diesen werden die Frequenzgänge beeinflusst, die Ein- und Ausgänge gleichspannungsfrei gemacht und eine unerwünschte Selbsterregung über die Betriebsspannungszuführung verhindert.

Zusammenfassung

In Niederfrequenzverstärkern müssen häufig die Verstärkung (Lautstärke und Balance) und der Frequenzverlauf (Klangbild) einstellbar sein. Man benötigt dazu verschiedene Einstellglieder.

Als Verstärkungseinsteller ist vor dem Verstärkereingang ein stetig veränderbarer Spannungsteiler vorgesehen, der zur Anpassung an die Empfindlichkeitskurve des Ohres sein Teilerverhältnis logarithmisch mit dem Drehwinkel verändern muß.

Um eine gehörrichtige (physiologische) Lautstärkeeinstellung zu realisieren, muß das verwendete Potentiometer einen Abgriff haben und mit einer RC-Schaltung kombiniert werden.

Zur Klangeinstellung muß ein umfangreicheres Netzwerk zwischen zwei Vorverstärkerstufen gelegt werden. Durch zwei Potentiometer, die als Höhen- und als Tiefeneinsteller dienen, läßt sich eine mehr oder weniger starke Bevorzugung der hohen und tiefen Frequenzen erreichen.

Elektronische Einstellglieder sind aufwendiger; sie haben aber die Vorteile, daß keine abgeschirmten Leitungen zu den Einstellorganen geführt werden müssen, keine mechanisch gekoppelten Potentiometer benötigt werden und daß der Bedienungskomfort der Geräte erhöht werden kann. Elektronische Einstellglieder werden meistens mit speziellen integrierten Schaltkreisen realisiert, die kreuzgekoppelte Differenzverstärker enthalten. Bei diesen läßt sich die Stromverteilung in den einzelnen Zweigen steuern und damit die Verstärkung und der Frequenzgang.

3.8. Aufgaben

1. Erläutern Sie die grundsätzlichen Verstärkerprobleme beim Vor-, End- und Senderverstärker!
2. Geben Sie die Verstärkungen $v_p = 1000$; $v_u = 100$; $v_i = 10$ in dB an!
3. Erläutern Sie die physikalische Wirkungsweise der Verstärkergrundsaltungen mit einem Transistor!
4. Erläutern Sie, warum bei der Emitterschaltung eine Gegenphasigkeit von Ausgangs- zu Eingangsspannung auftritt!
5. Begründen Sie das Absinken des Eingangswiderstands bei der Basisschaltung und das Ansteigen desselben bei der Kollektorschaltung gegenüber der Emitterschaltung!
6. Führen Sie eine zu Aufgabe 5 analoge Betrachtung für den Ausgangswiderstand durch!
7. Erläutern Sie, warum durch das Bootstrap-Prinzip der Eingangswiderstand eines Verstärkers dynamisch vergrößert werden kann!

8. Begründen Sie die Zonenfolge des Ersatztransistors der
 - a) Standard-Darlington-Schaltung
 - b) Komplementär-Darlington-Schaltung!
9. Stellen Sie die Vor- und Nachteile gegenüber, die sich bei der Kombination eines Unipolar- mit einem Bipolartransistor ergeben!
10. Erläutern Sie, warum bei der Differenzverstärkerschaltung die Ausgangsspannungsänderung proportional der Differenz der beiden Eingangsspannungen ist!
11. Erläutern Sie die Wirkungsweise der Stromquellschaltung nach Bild 3.12!
12. Nennen Sie die Vorteile einer mit Komplementärtransistoren aufgebauten Koppelschaltung gegenüber einer mit Transistoren gleichen Leitungstyps aufgebauten Koppelschaltung!
13. Nennen Sie die Vor- und Nachteile, die ein gegengekoppelter Verstärker gegenüber einem nichtgegengekoppelten hat!
14. Begründen Sie, warum der Betrag des Rückkopplungsgrads größer als 1 sein muß, damit eine Gegenkopplung auftritt!
15. Erläutern Sie, wie sich bei den vier Gegenkopplungsgrundschaltungen die Eingangs- und Ausgangswiderstände mit zunehmendem Rückkopplungsgrad verändern!
16. Nennen und erläutern Sie die verschiedenen Kenngrößen von Operationsverstärkern!
17. Erläutern Sie, wie beim integrierten Operationsverstärker durch eine äußere Beschaltung eine Frequenzkompensation, Offsetkompensation und Driftminderung erreicht werden kann!
18. Entwerfen Sie die komplette Schaltung eines invertierenden Verstärkers mit einem OPV für $|v| = 40 \text{ dB}$ und $R_{\text{ein}} = 10 \text{ k}\Omega$!
19. Erläutern Sie die Wirkungsweise des Operationsverstärkers mit Stromdifferenzeingang!
20. Erläutern Sie den typischen Frequenzverlauf eines RC-Verstärkers!
21. Geben Sie an, durch welche Maßnahmen man beim RC-Verstärker die Bandbreite vergrößern kann!
22. Erläutern Sie anhand des Stromlaufplans die physikalische Wirkungsweise des ZF-Verstärkers eines industriell gefertigten Rundfunkempfängers!
23. Begründen Sie, warum bei direktgekoppelten Verstärkern eine Wechsellspannungsgegenkopplung über mehrere Stufen leichter realisierbar ist als bei RC-gekoppelten!
24. Erläutern Sie die Wirkungsweise des direktgekoppelten Verstärkers nach Bild 3.63!
25. Erläutern Sie die Vorteile eines Gegentaktverstärkers gegenüber einem Eintaktverstärker!
26. Stellen Sie die Vor- und Nachteile des Gegentakt-AB- und Gegentakt-B-Verstärkers gegenüber!
27. Entwerfen Sie die Schaltung eines quasikomplementären Gegentaktverstärkers!
28. Entwerfen Sie die Schaltung eines mit einem integrierten NF-Verstärker angesteuerten komplementären Gegentaktverstärkers!
29. Erläutern Sie das Prinzip einer elektronischen Ausgangsstrombegrenzung!
30. Erläutern Sie anhand des Übersichtsschaltplans die prinzipielle Wirkungsweise des D-Verstärkers!
31. Erläutern Sie den Zweck und die Wirkungsweise der gehörrichtigen (physiologischen) Lautstärkeeinstellung!
32. Erläutern Sie die Wirkungsweise des Klangeinstellnetzwerkes nach Bild 3.100!
33. Nennen Sie die Vor- und Nachteile elektronischer gegenüber nichtelektronischer Einstellglieder!
34. Erläutern Sie das Prinzip der Stromverteilungssteuerung bei einem Differenzverstärker!

4. Sinus- und Kippgeneratoren

4.1. Allgemeines

Unter Generatoren werden hier solche Schaltungen verstanden, die aus einer Gleichspannung eine Wechselspannung erzeugen. Diese Wechselspannung kann *sinusförmig* oder *nicht-sinusförmig* sein. Für diese Generatoren ist auch die Bezeichnung „Oszillator“ gebräuchlich. Zur Anregung einer Schwingung wird ein *verstärkendes Bauelement* oder eins mit *negativem Widerstand* benötigt. Hier sollen nur die Verhältnisse an einem verstärkenden Vierpol mit Rückkopplung untersucht werden (s. Abschn. 3.3.). Bild 4.1 zeigt eine entsprechende Darstellung. Besteht zwischen dem Teil der zurückgeführten Ausgangsspannung und der für

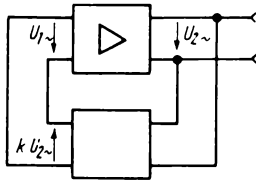


Bild 4.1. Vierpoldarstellung der Mitkopplung

diese Ausgangsspannung erforderlichen Eingangsspannung Gleichphasigkeit, so liegt eine Mitkopplung vor. Die für die Ausgangsspannung $U_{2\sim}$ benötigte Eingangsspannung $U_{1\sim}$ wird nach

$$U_{1\sim} = \frac{U_{2\sim}}{v_u} \quad (4.1)$$

berechnet. Wird dabei

$$kU_{2\sim} = U_{1\sim},$$

so ist keine äußere Eingangsspannung mehr erforderlich. Die Schaltung erregt sich selbst. Aus obiger Gleichung folgt

$$k = \frac{1}{v_u} \quad \text{oder} \quad kv_u = 1 \quad (4.2)$$

Der Ausdruck $kv_u = 1$ wird deshalb als *Selbsterregungsbedingung* bezeichnet. Ist sie nur für eine

Frequenz erfüllt, so werden Sinusschwingungen entstehen (s. Abschn. 4.2.); ist sie hingegen frequenzunabhängig erfüllt, so entstehen nicht-sinusförmige Ausgangsspannungen (s. Abschn. 4.3.).

Tritt am Verstärkervierpol eine Phasendrehung von 180° auf ($v_u < 0$), so muß der Mitkopplungsvierpol ebenfalls phasendrehend wirken ($k < 0$). Dieser Fall liegt beispielsweise bei Verwendung eines Transistors in Emitterschaltung vor. Wirkt der Verstärkervierpol nicht phasendrehend, darf auch der Mitkopplungsvierpol keine Phasendrehung aufweisen (z. B. Transistor in Basisschaltung). Für $kv_u < 1$ ist keine Selbsterregung möglich. Bei $kv_u > 1$ wird die Amplitude der angeregten Schwingung bis zur Selbstbegrenzung steigen.

Da $kv_u = 1$ einen labilen Grenzfall darstellt, wählt man in der Praxis $kv_u > 1$. Um trotzdem eine stabile Amplitude zu erhalten, ist eine *Amplitudenbegrenzung* erforderlich.

Eine Möglichkeit der Begrenzung ist die *Sättigungsbegrenzung*. Sie wird vielfach bei einfachen Transistoroszillatoren angewendet. Bild 4.2

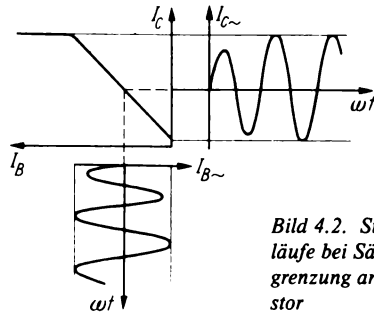


Bild 4.2. Stromverläufe bei Sättigungsbegrenzung am Transistor

zeigt die dynamische Stromverstärkungskennlinie eines npn-Transistors. Der Arbeitspunkt ist in üblicher Form festgelegt. Der obere Knick der Kennlinie entsteht durch die Begrenzung des Kollektorwechselstroms auf einen Maximalwert durch den Arbeitswiderstand. Bild 4.2 verdeutlicht, wie der Kollektorwechselstrom und damit auch der Basiswechselstrom begrenzt sind.

Nach dem beschriebenen Vorgang beeinflusst auch der nichtlineare Eingangswiderstand des Transistors die Begrenzung. Beim npn-Transistor wird stets die positive Halbwelle der Ausgangsspannung stärker durch den Transistoreingang belastet als die negative Halbwelle. In einigen Schaltungen umgeht man diesen Einfluß durch Anwendung der im Abschn. 4.2.1. beschriebenen Stromsteuerung.

Für hochwertige Oszillatoren gibt es verschiedene Möglichkeiten der *äußeren Amplitudenbegrenzung*. Eine der Formen wird im Bild 4.3 gezeigt. Die Antiparallelschaltung der Dioden begrenzt beide Halbwellen der Wechselspannung am Transistoreingang. Bild 4.4 zeigt die resultierende Kennlinie der antiparallelgeschalteten Dioden mit eingetragener begrenzter Wechselspannung.

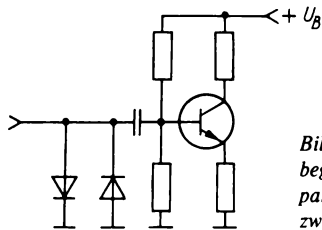


Bild 4.3. Amplitudenbegrenzung durch Antiparallelschaltung von zwei Dioden

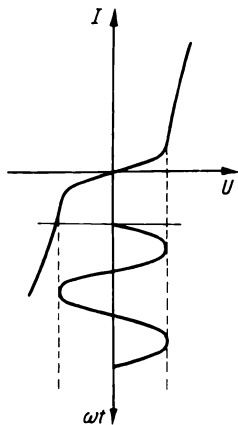


Bild 4.4. Darstellung der Begrenzung an der resultierenden Kennlinie der beiden Dioden im Bild 4.3

Vielfach werden an Oszillatoren hohe Anforderungen an die *Frequenzkonstanz* gestellt. Sie hängt sowohl von der Konstanz der Werte der frequenzbestimmenden passiven Bauelemente (beeinflusst durch Temperaturschwankungen und Alterungserscheinungen) als auch von Parameteränderungen der aktiven Bauelemente ab (hervorgerufen durch Betriebsspannungsschwankungen, Temperaturänderungen und Alterungserscheinungen). Bei geforderter hoher

Frequenzkonstanz sind folglich die genannten Einflußgrößen nach Möglichkeit unwirksam zu machen. Als Maßnahmen kommen u. a. eine Kompensation der TK-Werte, eine eventuelle Unterbringung der Schaltung in einem Thermostat, eine Regelung der Betriebsspannung, spezielle Schaltungsmaßnahmen zur Verringerung der Einflußgrößen und eine Voralterung (künstliche Alterung) der Bauelemente in Betracht.

4.2.

Sinusgeneratoren

4.2.1.

Meißner-Oszillatoren

Einige spezielle Oszillatorschaltungen werden im folgenden auf die Erfüllung der im Abschnitt 4.1. genannten Bedingungen untersucht. Die erforderliche Frequenzselektivität erreicht man entweder mittels eines frequenzabhängigen Arbeitswiderstands oder eines frequenzabhängigen Mitkopplungsvierpols. Beim *Meißner-Oszillator* findet ein Schwingkreis Anwendung. Der Transistor ist in Emitterschaltung geschaltet ($\Delta\varphi = 180^\circ$); deshalb ist im Mitkopplungsvierpol ebenfalls eine Phasendrehung von 180° erforderlich. Sie wird mit einem Transformator erreicht.

Eine entsprechende Schaltung wird im Bild 4.5 gezeigt. Der Arbeitspunkt ist über einen Basisspannungsteiler eingestellt und durch den Emittewiderstand stabilisiert. Der Spannungsteiler ist an das „kalte Ende“ der Mitkoppelspule gelegt, um die Belastung des Schwingkreises möglichst niedrig zu halten. Der Kondensator sorgt dafür, daß die Mitkoppelspule mit ihrem „kalten Ende“ wechselstrommäßig auf Masse liegt. Ein Nachteil dieser Schaltung besteht darin, daß der Kollektorgleichstrom die Schwingkreis-

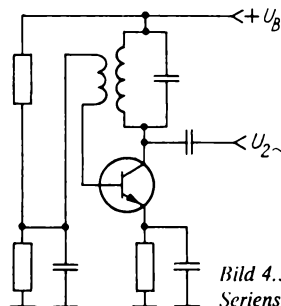


Bild 4.5. Meißner-Oszillator in Serienspeisung

induktivität durchfließt (Serienspeisung). Geringe Betriebsspannungsschwankungen ändern bei Spulen mit Kern die Vormagnetisierung, damit die Induktivität und Resonanzfrequenz des Schwingkreises.

Bei Anwendung der Parallelspeisung entfällt der Nachteil der Vormagnetisierung (Bild 4.6). Der Wechselstromarbeitswiderstand des Transistors wird durch die Parallelschaltung des Schwingkreises mit der Drossel bestimmt. Dabei soll der induktive Widerstand der Drossel wesentlich größer als der Resonanzwiderstand des Schwingkreises sein. Die Drossel soll dabei einen Kurzschluß der am Kollektor liegenden Wechselspannung über das Netzteil verhindern. Der dem Schwingkreis vorgeschaltete Kondensator dient der Gleichspannungstrennung. Der mit der Mitkopplerspule in Reihe liegende Widerstand führt zur *Stromsteuerung* des Transistors. Hierdurch ergibt sich eine Erniedrigung des Klirrfaktors. Die Amplitudenbegrenzung erfolgt bei fest eingestelltem und durch die Emittterkombination stabilisiertem Arbeitspunkt durch Sättigung. Dabei wird der Sättigungswechselstrom durch den Wechselstromlastwiderstand bestimmt. Das Potential am Kollektor entspricht der Betriebsspannung.

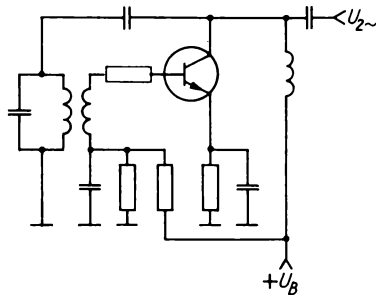


Bild 4.6. Meißner-Oszillator mit Parallelspeisung und Stromsteuerung

4.2.2. Dreipunktoszillatoren

Bei Dreipunktoszillatoren sorgt wiederum ein Schwingkreis dafür, daß die Selbsterregungsbedingung nur für eine Frequenz erfüllt ist. Im Schwingkreis wird entweder die Spule mit einer Mittelanzapfung versehen oder der Kondensator als Spannungsteiler ausgeführt, wobei jeweils der Mittelpunkt an Masse liegt. Bild 4.7 zeigt eine kapazitive Dreipunktschaltung in Parallelspeisung (auch *Colpitts-Oszillator* ge-

nannt). Im Bild 4.8 ist der Schwingkreis so umgezeichnet, daß Masse als Bezugspunkt für die eingetragenen Spannungen dient. Das zugehörige Zeigerbild soll die Phasendrehung zwischen Kollektor- und Basisspannung verdeutlichen. Die Konstruktion des Zeigerbilds beginnt z. B. mit der Basiswechselspannung $U_{1\sim}$. $I_{1\sim}$ muß $U_{1\sim}$ um 90° voreilen; ebenso muß $U_{L\sim}$ dem Strom $I_{1\sim}$ um 90° voreilen. Ist nun $U_{L\sim}$ größer als $U_{1\sim}$ (wie im Beispiel angenommen), so ist die vektorielle Summe aus $U_{L\sim}$ und $U_{1\sim}$, die gleich der Kollektorspannung $U_{2\sim}$ ist, gegenüber der Basiswechselspannung um 180° pha-

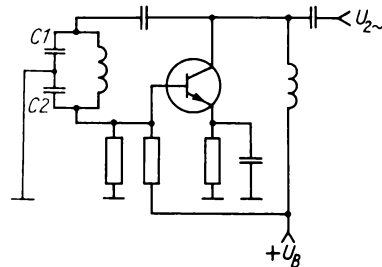


Bild 4.7. Kapazitive Dreipunktschaltung

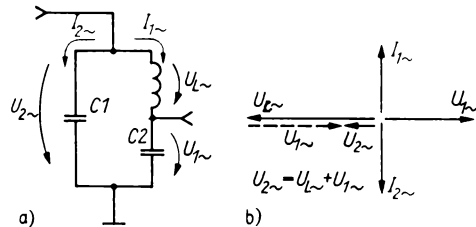


Bild 4.8. Phasendrehung bei kapazitiver Dreipunktschaltung

a) Schwingkreis; b) Zeigerdiagramm

senverschoben. Der Mitkopplungsfaktor ist durch die beiden Kondensatoren bestimmt und wird

$$k = \frac{C_1}{C_2} \quad (4.3)$$

Die Colpitts-Schaltung weist eine geringe Abhängigkeit der Betriebsfrequenz von der arbeitspunktabhängigen Eingangs- und Ausgangskapazität des Transistors auf. Diese liegen hier parallel zu den Kondensatoren $C1$ und $C2$, die höhere Kapazitäten haben als die wirksame Schwingkreiskapazität (Reihenschaltung von $C1$ und $C2$).

Im Bild 4.9 ist eine induktive Dreipunktschaltung dargestellt (auch *Hartley-Oszillator* genannt). Der Mittelpunkt des induktiven Spannungsteilers liegt hier an der Betriebsspannung und damit ebenfalls wechselstrommäßig an Masse. (Die Betriebsspannung hat stets wechselstrommäßig Massepotential.)

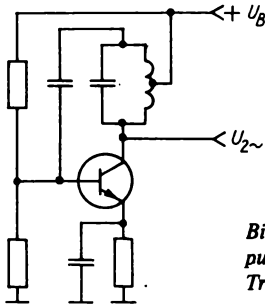


Bild 4.9. Induktive Dreipunktschaltung mit einem Transistor

Will man eine verstellbare Frequenz an einem LC-Oszillator erreichen, kann man entweder die Kapazität C (Drehkondensator) oder die Induktivität L (Variometer) verstellbar ausführen. In beiden Fällen handelt es sich um mechanisch komplizierte und damit auch störanfällige und teure Bauelemente. In verschiedenen Anwendungsfällen ist es möglich, eine Kapazitätsdiode in den Schwingkreis einzubauen, deren wirksame Kapazität mit der Höhe der angelegten Sperrspannung beeinflusst wird. Derartige spannungsgesteuerte Oszillatoren werden auch mit der Kurzbezeichnung VCO bezeichnet (engl., voltage controlled oscillator). Im Bild 4.10 ist ein VCO in Anlehnung an die Oszillatorschaltung nach Bild 4.6 dargestellt. Die Kapazitätsdiode wird durch E in Sperrrichtung vorgespannt. C verhindert einen Kurzschluß der Sperrspannung der Kapazitätsdiode. Die Drossel L ist erforderlich, damit die Ausgangswechselspannung nicht über die die Steuerspannung

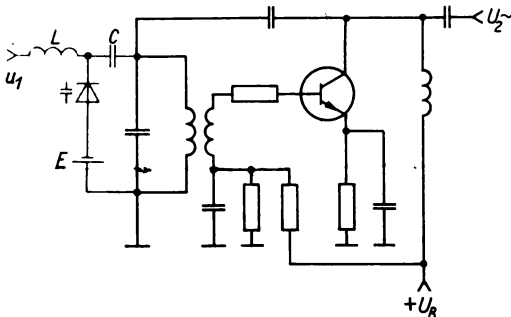


Bild 4.10. Spannungsgesteuerter Oszillator (VCO)

u liefernde Quelle kurzgeschlossen wird. Durch Änderung der Spannung u ändert sich die wirksame Sperrspannung an der Kapazitätsdiode, damit ändern sich die Kapazität der Kapazitätsdiode und die Frequenz des Oszillators.

4.2.3.

Hochfrequenzoszillator in Basisschaltung

Für höhere Frequenzen setzt man bei Transistoroszillatoren die Basisschaltung ein. Diese verursacht keine Phasendrehung. Deshalb darf auch der Mitkopplungsvierpol keine Phasendrehung verursachen. Bild 4.11 zeigt eine entsprechende Schaltung. Der als Arbeitswiderstand dienende Schwingkreis sorgt auch hier dafür, daß v_u nur für die Resonanzfrequenz eine ausreichende Größe erhält. Die Mitkopplung erfolgt durch Abgriff eines Teiles der Ausgangsspannung von der Schwingkreisinduktivität. Dabei entsteht gleichzeitig eine Widerstandstransformation zur Anpassung des niederohmigen Eingangswiderstands an den hochohmigen Ausgangswiderstand der Basisschaltung. Eine abermalige Teilung der Mitkopplungsspannung erfolgt über die Kapazität und den in Reihe liegenden Eingangswiderstand. Der Kondensator zwischen Basis und Masse soll das erforderliche wechsellspannungsmäßige Massepotential der Basis schaffen. Die in der Emittierleitung liegende Drossel verhindert den Kurzschluß der am Emittier liegenden Wechselspannung. Der Arbeitspunkt wird durch den Basisspannungsteiler bestimmt.

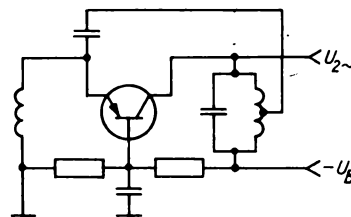


Bild 4.11. Hochfrequenzoszillator in Basisschaltung

4.2.4.

Quarzoszillatoren

Für viele Anwendungsfälle ist die mit den vorstehend beschriebenen Oszillatoren erreichbare Frequenzstabilität nicht ausreichend. Oszillatoren mit Frequenzabweichungen von 10^{-5} bis 10^{-9} sind mit Schwingquarzen realisierbar. Da-

bei nutzt man, daß Schwingquarze wie Schwingkreise sehr hoher Güte ($10^4 \dots 10^7$) wirken. Die meist zusätzlich verwendeten LC-Schwingkreise sollen Störschwingungen auf Nebenwellen des Quarzes vermeiden. In geringem Umfang ist ein Ziehen der Frequenz möglich (Veränderung der Schwingfrequenz durch Zuschalten von Blindwiderständen). Den Quarz kann man sowohl in Serienresonanz als auch in Parallelresonanz betreiben.

Ein Beispiel für Parallelresonanz ist im Bild 4.12 dargestellt. Es handelt sich um eine induktive Dreipunktschaltung. Bei Abweichung von der Resonanzfrequenz stellt der im Nebenschluß zum Transistoreingang liegende Quarz einen niedrigen Widerstand dar – eine Selbsterregung ist nicht möglich. Zur Erhöhung des Eingangswiderstands ist der Transistor gegengekoppelt.

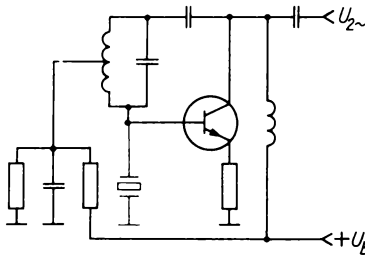


Bild 4.12. Quarzoszillator in Parallelresonanz

Im Bild 4.13 ist ein Quarzoszillator in Basis-schaltung dargestellt. Der Quarz wird in Serienresonanz betrieben und ergibt nur für seine Serienresonanzfrequenz eine genügend große Mitkopplung.

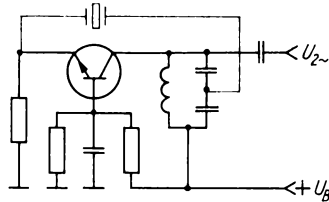


Bild 4.13. Quarzoszillator in Serienresonanz

Wird die Möglichkeit des Ziehens der Frequenz durch Zuschaltung einer Kapazitätsdiode genutzt, spricht man von einem VCXO (engl., voltage controlled xtal oscillator). Entsprechend den typischen Eigenschaften des Quarzoszillators ist die erreichbare relative Frequenzänderung durch Änderung der Sperrspannung der Kapazitätsdiode sehr gering.

4.2.5. RC-Oszillatoren

In den bisherigen Schaltungen wurde die Selbsterregungsbedingung durch Einsatz eines Schwingkreises nur für eine Frequenz erfüllt. Für niedrige Frequenzen wird die dabei erforderliche Induktivität sehr groß. Einfacher sind hier Schaltungen mit RC-Netzwerken. Zwei Formen dieser Netzwerke finden für Oszillatoren häufig Anwendung: die *Phasenschieberkette* und die *Wien-Brücke*. Einen RC-Oszillator mit Phasenschieberkette zeigt Bild 4.14., Im Mit-

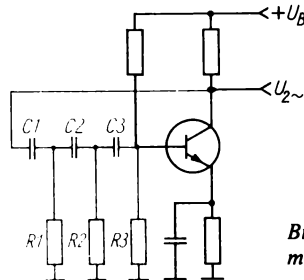


Bild 4.14. RC-Oszillator mit Phasenschieberkette

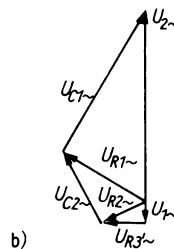
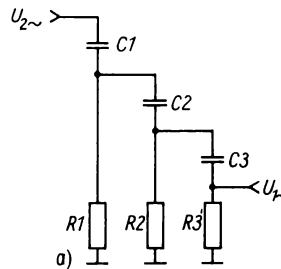


Bild 4.15.
RC-Phasenschieberkette
a) Schaltung;
b) Zeigerbild, vereinfacht

kopplungsvierpol muß hierbei sowohl die Phasendrehung als auch eine ausreichende Größe von k nur für eine Frequenz erreicht werden. Bild 4.15a zeigt den Mitkopplungsvierpol einzeln. Es handelt sich um die Hintereinanderschaltung von drei Wechselspannungsteilern. Zur Vereinfachung wird jeder Spannungsteiler als unbelastet betrachtet. Die anliegende Span-

nung wird jeweils in zwei rechtwinklig zueinander liegende Teilspannungen aufgeteilt (Bild 4.15b). Die Spannung $U_{1\sim}$ weist nur für eine bestimmte Frequenz eine Phasendrehung von 180° auf, da das Verhältnis der Katheten der einzelnen Dreiecke frequenzabhängig ist. (Stiege die Frequenz, so würden die kapazitiven Spannungsabfälle kleiner, und die erreichte Phasendrehung wäre $< 180^\circ$.)

Entsprechende RC -Ketten werden sowohl mit drei als auch mit vier oder fünf RC -Gliedern ausgeführt. Zur angenäherten Bestimmung der Frequenz bei einer dreigliedrigen RC -Kette dient die Gleichung

$$f = \frac{1}{2\pi\sqrt{6}CR} \quad (4.4)$$

Dabei ist vorausgesetzt, daß alle Kondensatoren die gleiche Kapazität und alle Widerstände den gleichen Widerstandswert haben. Aus diesem Grunde ist $R/3$ so zu bemessen, daß der Widerstandswert seiner Parallelschaltung mit dem Transistoreingangswiderstand und dem oberen Basisspannungsteilerwiderstand gleich dem Wert von $R/1$ und $R/2$ wird.

Einen einfachen *Wien-Brücken-Oszillator* zeigt Bild 4.16. Auch hier wird zunächst der Rückkopplungszweig näher betrachtet. Aus der Darstellung im Bild 4.17a folgt das Zeigerbild 4.17b. Man erkennt, daß für eine bestimmte Frequenz die Phasenverschiebung zwischen Eingang und Ausgang Null wird. Die Konstruktion des Zeigerbildes kann man mit $U_{1\sim}$ beginnen. Aus dieser Spannung folgen die beiden Ströme $I_{R\sim}$ und $I_{C\sim}$, deren Summe $I_{\sim}$ ergibt. Dieser Strom verursacht an der Reihenschaltung aus R und C entsprechende Spannungsabfälle. Die Summe der Wechselspannungen $U_{R\sim}$, $U_{C\sim}$ und $U_{1\sim}$ ergibt die Spannung $U_{2\sim}$. Bei Änderung der Frequenz ergibt sich zwischen $U_{1\sim}$ und $U_{2\sim}$ eine Phasenverschiebung. (So wird bei Frequenzerniedrigung $I_{R\sim} > I_{C\sim}$ und $U_{C\sim} > U_{R\sim}$; daraus folgt eine Voreilung von $U_{1\sim}$ gegenüber $U_{2\sim}$.) Da der obere Teil der Schaltung einen Hochpaß und der untere Teil einen Tiefpaß darstellt, folgt der Amplitudengang der Schaltung nach Bild 4.17c.

Das betrachtete Netzwerk verursacht keine Phasendrehung; es ist bei der Verwendung der Emitterschaltung eine weitere Stufe zur Phasendrehung erforderlich. Die Mindestverstärkung beider Stufen beträgt nur 3, weshalb eine starke Gegenkopplung notwendig ist. Über die beiden

antiparallelgeschalteten Dioden wird eine zusätzliche amplitudenabhängige Gegenkopplung erreicht und damit eine konstante Amplitude der Ausgangswechselspannung.

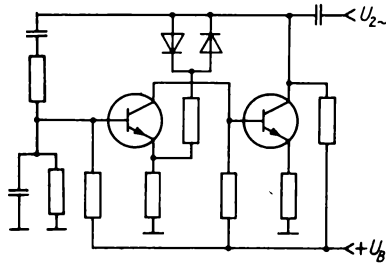
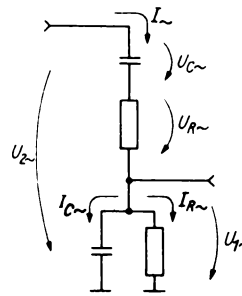
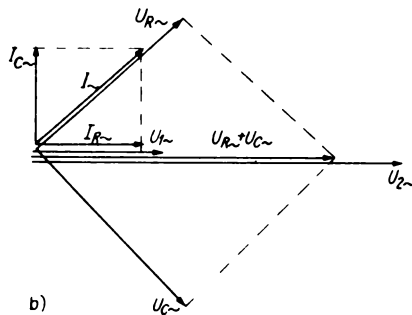


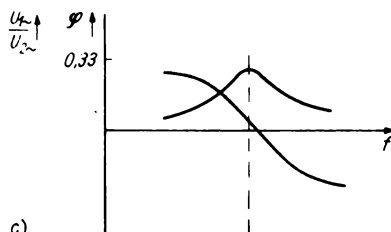
Bild 4.16. Wien-Brücken-Oszillator



a)



b)



c)

Bild 4.17. Wien-Brücke

a) Schaltung; b) Zeigerbild; c) Phasen- und Amplitudenfrequenzgang

Wien-Brücken-Oszillatoren eignen sich besonders für durchstimmbare NF-Generatoren, wobei beide Widerstände oder beide Kondensatoren gleichzeitig verändert werden müssen. Die Schwingfrequenz ergibt sich zu

$$f = \frac{1}{2 \pi CR} \quad (4.5)$$

Auch diese Gleichung gilt nur für gleiche Kapazität beider Kondensatoren und gleiche Widerstandswerte beider Widerstände.

Bei RC-Oszillatoren hängt die Frequenzänderung direkt von der R - oder C -Änderung ab [s. Gl. (4.4) und (4.5)], bei LC-Oszillatoren nur von der Wurzel der Kapazitätsänderung. Ist beispielsweise ein Drehkondensator mit einem V_{Δ} variationsbereich 1:9 vorhanden, so ist beim RC-Oszillator eine Frequenzänderung von 1:9 und beim LC-Oszillator eine von 1:3 erreichbar.

4.2.6.

Oszillatoren mit integrierten Schaltkreisen

Auch mit integrierten Schaltkreisen sind Oszillatoren aufbaubar. Dabei gelten grundsätzlich die im Abschn. 4.1. genannten Voraussetzungen. Neben speziellen Schaltkreisen eignet sich der Operationsverstärker gut als verstärkendes Element in Oszillatorschaltungen.

Zuerst soll eine Oszillatorschaltung mit dem Schaltkreis A 301 betrachtet werden. Der Schaltungsausschnitt der IS zeigt einen von einer intern stabilisierten Spannung gespeisten mehrstufigen Verstärker (Bild 4.18). Über den äußeren Widerstand erfolgt eine frequenzunabhängige Mitkopplung. (Die Phasenverhältnisse sind für die positive Halbwelle an der Basis von $V5$ rot eingetragen.) Der Schwingkreis liegt im Gegenkopplungszweig und sorgt dafür, daß nur bei seiner Resonanzfrequenz die Gegenkopplung so niedrig wird, daß eine Selbsterregung eintreten kann. Folglich wird eine Sinusspannung entstehen. Der Transistor $V4$ gehört zur internen Spannungsstabilisierung. $V7$ sorgt als Emitterfolger für eine geringe Belastung der Phasenumkehrstufe $V6$. Über die Schwingkreisspule erhält die Basis von $V5$ das Emitterpotential von $V6$. Zur Erklärung sei angenommen, daß das Potential am Emitter von $V6$ (und damit an der Basis von $V5$) 1,0 V betrage. Es folgt daraus für die Basis von $V6$ und den Kollektor von $V5$ ein um eine Diodenflußspannung höheres Potential

($\approx 1,7$ V) und für den Emitter von $V5$ entsprechend $\approx 0,3$ V. Der Spannungsabfall über $V5$ von $\approx 1,4$ V ergibt eine ausreichende Aussteuerbarkeit.

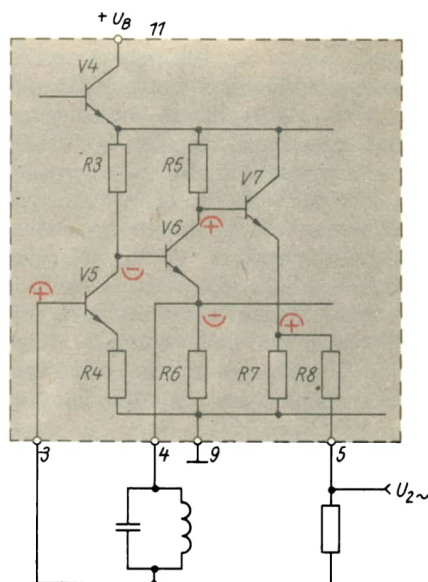


Bild 4.18. Oszillator mit dem integrierten Schaltkreis A 301

Einen einfachen Wien-Brücken-Oszillator mit Operationsverstärker zeigt Bild 4.19. (Frequenz- und Offsetkompensation sind nicht mit dargestellt.) Im Abschn. 4.2.5. wurde bereits ein Wien-Brücken-Oszillator erläutert. Da das Netzwerk keine Phasendrehung aufweist, wird es an den nicht invertierenden Eingang des Operationsverstärkers angeschlossen. Die mitgekoppelte Spannung beträgt etwa $\frac{1}{3}$ der Ausgangsspannung; die Verstärkung des integrierten Schaltkreises ist entsprechend niedrig einzustellen. Dazu dient das Gegenkopplungsnetz-

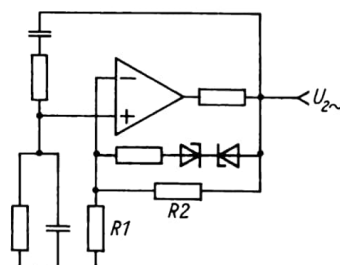


Bild 4.19. Wien-Brücken-Oszillator mit Operationsverstärker

werk am invertierenden Eingang. Die durch den ohmschen Teiler ($R1$, $R2$) festgelegte Verstärkung soll dabei etwas größer als 3 sein. Einem unzulässigen Ansteigen der Amplitude wird durch die beiden Z-Dioden entgegengewirkt (im Bild 4.19 rot eingetragen). Steigt die Ausgangsspannung, so wird der Widerstand der Z-Dioden zunehmend kleiner; damit sinkt die Verstärkung. Der ursprünglichen Amplitudenänderung wird entgegengewirkt. Um bei dieser Schaltung den Klirrfaktor niedrig zu halten, ist zu den Z-Dioden ein Widerstand in Reihe geschaltet, der einen Ausgleich des scharfen Kennlinienknicks bewirkt.

Einen Wien-Brücken-Oszillator mit Amplitudenbegrenzung durch einen als steuerbaren Widerstand verwendeten MOS-Transistor zeigt Bild 4.20 (ohne Frequenzgang- und Offsetkompensation dargestellt). Diese Schaltung hat den Vorteil eines niedrigen Klirrfaktors. Die Ausgangsspannung $U_{2\sim}$ wird gleichgerichtet und durch ein RC-Siebglied (Tiefpaß) geglättet. Die der Spannung $U_{2\sim}$ proportionale Gleichspannung am Kondensator wird dem Gate des MOS-Transistors zugeführt. Die Verstärkung des Operationsverstärkers wird durch die Widerstände $R1$ und $R2$ auf kleiner als 3 eingestellt. Die nochmalige Teilung der gegengekoppelten Spannung ermöglicht eine von dieser Teilung abhängige Erhöhung der Verstärkung auf den gewünschten Wert von $v_u = 3$. Stiege nun die Ausgangsspannung, so würde die dann auch negativer werdende Gatespannung eine Erhöhung des Widerstandes des Transistors bewirken. Eine Erniedrigung der Spannungsverstärkung des Operationsverstärkers wäre die Folge. Die Verwendung des Potentiometers gestattet zusätzlich zur Amplitudenbegrenzung eine Einstellung der Amplitude der Ausgangsspannung.

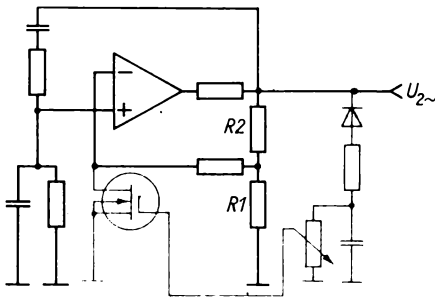


Bild 4.20. Wien-Brücken-Oszillator mit geringem Klirrfaktor

Zusammenfassung

Ein mitgekoppelter Verstärker erzeugt bei Einhaltung der Selbsterregungsbedingung Schwingungen. Die Schaltung darf zur Anregung einer Sinusschwingung die Mitkopplungsbedingung nur für eine Frequenz erfüllen. Dabei ist es gleichgültig, ob v_u oder k oder beide frequenzabhängig sind. Wird der aktive Vierpol in einer phasendrehenden Schaltung betrieben, so muß der Mitkopplungsvierpol ebenfalls phasendrehend wirken. Eine stabile Amplitude erreicht man durch Amplitudenbegrenzung.

Beim Meißner-Oszillator wird die erforderliche Phasendrehung durch einen Transformator hervorgerufen. Die Frequenz bestimmt ein Schwingkreis. Die Stromversorgung erfolgt durch Serien- oder Parallelspeisung.

Bei Dreipunktoszillatoren wird die erforderliche Phasendrehung durch einen Spannungsteiler mit Mittelabgriff hervorgerufen. Ein Schwingkreis sorgt dafür, daß die Mitkopplungsbedingung nur für eine Frequenz erfüllt ist und somit eine Sinusspannung entsteht.

Bei einem Oszillator in Basisschaltung ist keine Phasendrehung im Mitkopplungsvierpol erforderlich. Die Frequenzbestimmung erfolgt durch einen Schwingkreis. Zur Widerstandsanpassung des hochohmigen Ausgangswiderstands an den niederohmigen Eingangswiderstand der Basischaltung erfolgt die Mitkopplung über eine Anzapfung der Schwingkreisspule.

Eine hohe Frequenzkonstanz wird mit Quarzoszillatoren erreicht. Der Quarz kann dabei sowohl in Serienresonanz als auch in Parallelresonanz betrieben werden. Zusätzliche LC-Schwingkreise sollen eine Selbsterregung der Schaltung auf Nebenfrequenzen des Quarzes verhindern.

Die erforderliche Phasendrehung wird bei einem RC-Oszillator mit Phasenschieberkette durch eine Kettenschaltung von RC-Gliedern hervorgerufen. Gleichzeitig wird erreicht, daß nur bei einer Frequenz die Mitkopplungsbedingung erfüllt ist.

Beim Wien-Brücken-Oszillator erfolgt im Mitkopplungsvierpol keine Phasendrehung; der Ausgangsspannungsverlauf des Mitkopplungsvierpols trägt Bandpaßcharakter. Bei Verwendung der Emitterschaltung ist eine Phasenumkehrstufe erforderlich.

Oszillatoren mit integrierten Schaltkreisen werden nach den gleichen Gesichtspunkten wie Oszillatoren in diskreter Schaltungstechnik aufgebaut.

4.3. Kippgeneratoren

4.3.1. Allgemeines

Unter *Kippspannung* versteht man einen zeitlichen Spannungsverlauf mit Unstetigkeitsstellen (Knickpunkten). Unter anderem sind Dreieck-, Sägezahn- und Rechteckspannung Kippspannungen. In den meisten Fällen entstehen entsprechende Spannungsverläufe bei Oszillatorschaltungen infolge ständiger Umladung von Kondensatoren.

Bild 4.21 zeigt eine *RC*-Schaltung im Schaltbetrieb. Es wird angenommen, daß der Kondensator vor dem Zuschalten an die Spannungsquelle völlig entladen ist. Im Einschaltmoment beträgt nach

$$U_c = \frac{Q}{C} \quad (4.6)$$

die Spannung am Kondensator 0 V. Der Kondensator wird durch den fließenden Ladestrom nach der Beziehung

$$\Delta Q = I \Delta t \quad (4.7)$$

aufgeladen. Mit zunehmender Ladung sinkt der Ladestrom exponentiell, und die Spannung steigt am Kondensator. Diese Verläufe sind im Bild 4.22 dargestellt, wobei t_1 den Einschaltmoment festlegt. Nach vollendeter Aufladung entladen wir den Kondensator (Schalterstellung 2).

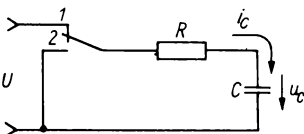


Bild 4.21. Schaltung zur Auf- und Entladung eines Kondensators

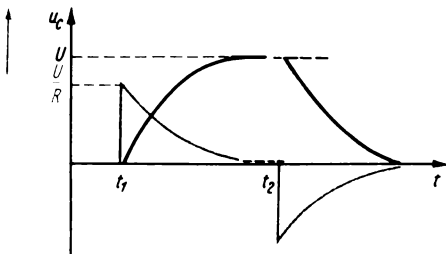


Bild 4.22. Strom- und Spannungsverlauf am Kondensator beim Auf- und Entladen

Der Beginn der Entladung ist im Bild 4.22 mit t_2 bestimmt. Der nun fließende Entladestrom hat die umgekehrte Richtung des Ladestroms. Entladestrom und Entladespannung sinken exponentiell auf den Wert Null.

An Kondensatoren gibt es keine sprunghaften Spannungsänderungen. Zum Vereinfachen der Berechnung derartiger Schaltvorgänge rechnet man mit der *Zeitkonstanten* τ .

τ gibt an, in welcher Zeit sich ein verlustloser Kondensator über einen parallelliegenden Widerstand auf den 0,368fachen Wert seiner Ladespannung entladen hat oder in welcher Zeit er sich über einen Widerstand auf den 0,632fachen Wert der Ladespannung aufgeladen hat.

Die Strom- und Spannungsverläufe werden durch e^x -Funktionen beschrieben. e ist die Basis der natürlichen Logarithmen ($e = 2,72$). Die gegebenen Zahlenwerte folgen aus der Beziehung

$$\frac{1}{e} = 0,368 \quad \text{und} \quad 1 - \frac{1}{e} = 0,632.$$

Die Auflade- und Entladezeitkonstante sind im Bild 4.23 eingetragen. Berechnet wird die Zeitkonstante einer *RC*-Schaltung nach

$$\tau = CR \quad (4.8)$$

Anstelle des Schalters im Bild 4.21 wird eine Rechteckspannungsquelle mit symmetrischem Ausgangsspannungsverlauf verwendet. Die Zeitdauer einer Halbperiode wird mit t_i bezeichnet und sei gleich der halben Periodendauer. Im folgenden wird der Ausgangsspannungsverlauf in Abhängigkeit von der Eingangsspannung für verschiedene Größenbeziehungen zwischen τ und t_i untersucht (Bild 4.24).

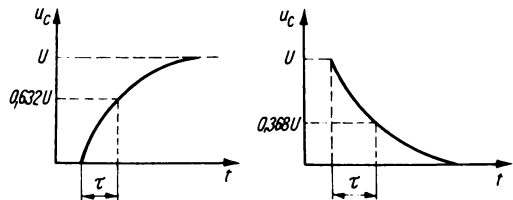


Bild 4.23. Auflade- und Entladezeitkonstante einer *RC*-Schaltung

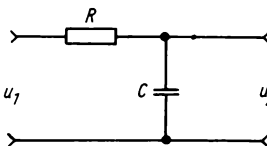


Bild 4.24. *RC*-Tiefpaß

Ist $\tau \ll t_i$, so kann der Kondensator in gegenüber t_i kurzer Zeit völlig umgeladen werden. Der Ausgangsspannungsverlauf nähert sich dem Eingangsspannungsverlauf. Die entstehende Abweichung wird mit Anstiegsverzögerung bezeichnet (Bild 4.25b).

Wird $\tau \approx t_i$, so erfolgt während der gesamten Impulsdauer t_i ein Aufladen und danach bis zum erneuten Beginn des Impulses ein Entladen des Kondensators mit typisch exponentiellem Verlauf (Bild 4.25c).

Für $\tau \gg t_i$ kann während der Impulsdauer t_i nur ein sehr kleiner Bereich der Aufladekurve und danach bis zum erneuten Beginn der Aufladung nur ein sehr kleiner Teil der Entladekurve durchlaufen werden. Dieser sehr kleine Ausschnitt ist dann nahezu linear (Bild 4.25d). Eine so dimensionierte RC-Schaltung wird als Integrierglied bezeichnet. Die Amplitude der Ausgangsspannung ist hier wesentlich kleiner als die Amplitude der Eingangsspannung.

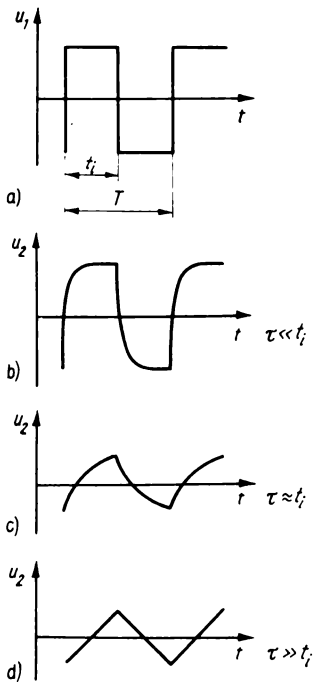


Bild 4.25. Ausgangsspannungsverlauf an einem Tiefpaß; Amplitudenverhältnisse nicht berücksichtigt

Nachfolgend wird das Schaltverhalten der Kombination nach Bild 4.26 untersucht. Diese Schaltung ist ähnlich der vorigen Schaltung; jedoch wird hier der Spannungsabfall am Wider-

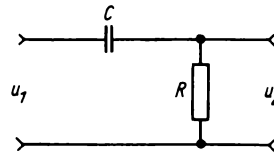


Bild 4.26.
RC-Hochpaß

stand und nicht der Spannungsabfall am Kondensator untersucht.

Ist hier $\tau \gg t_i$, so erfolgt während der einzelnen Impulse nur eine geringfügige Umladung des Kondensators. Die entstehende Veränderung der Form der Eingangsspannung bezeichnet man mit Dachabfall. Mit gegenüber t_i größer werdender Zeitkonstante τ nähert sich der Ausgangsspannungsverlauf immer mehr dem Eingangsspannungsverlauf (Bild 4.27b).

Wird $\tau \approx t_i$, so erfolgt eine stärkere Umladung des Kondensators über den Widerstand (Bild 4.27c).

Für $\tau \ll t_i$ wird der Kondensator während T ganz umgeladen. Vor Eintritt des positiven Spannungssprungs (plötzlicher Anstieg der Spannung auf einen positiven Wert) ist der Kondensator völlig entladen. Da es am Kondensator keine sprunghaften Spannungsänderungen gibt, muß ein positiver Spannungssprung am Eingang einen ebensolchen am Ausgang auslö-

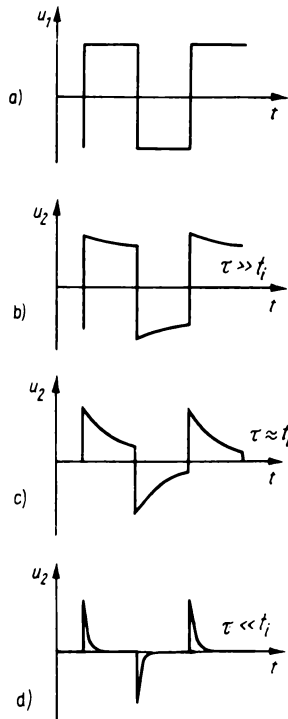


Bild 4.27. Ausgangsspannungsverlauf an einem Hochpaß

sen. Es erfolgt dann eine sehr rasche Aufladung des Kondensators auf den Wert der Eingangsspannung, wobei der Spannungsabfall am Widerstand auf 0 V sinkt. Gleiche Vorgänge entstehen beim negativen Spannungssprung am Eingang (Bild 4.27d). Eine so dimensionierte RC-Kombination wird als Differenzierglied bezeichnet.

Die Ausgangsspannungsverläufe der Schaltungen nach den Bildern 4.24 und 4.26 entstehen auch für unsymmetrische oder einer Gleichspannung überlagerte Rechteckspannungen im eingeschwungenen Zustand der Schaltung, wobei im Falle des Tiefpaß dann auch die Ausgangsspannung einer Gleichspannung überlagert ist.

4.3.2.

Rechteckspannungsgeneratoren

Der astabile Kippgenerator (auch Multivibrator genannt) stellt eine einfache Schaltung zur Erzeugung einer Rechteckspannung dar. Eine entsprechende Schaltung wird im Bild 4.28 gezeigt. (Signalgeneratoren in integrierter Schaltungstechnik werden im Abschn. 5.4.5. behandelt.)

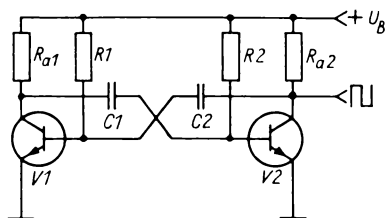


Bild 4.28. Astabiler Kippgenerator

Zur Erklärung der Wirkungsweise sei angenommen, daß Transistor V2 gerade völlig durchsteuert sei. Die negative Halbwelle der Kollektorspannung liegt über C2 an der Basis von V1 und sperrt diesen. Die damit am Kollektor von V1 liegende positive Halbwelle wird über C1 auf die Basis von V2 gekoppelt. Transistor V2 bleibt voll durchsteuert. Es entlädt sich der Kondensator C2 über den Widerstand R1 so lange, bis durch Transistor V1 ein geringer Kollektorstrom fließen kann. Die Spannung am Kollektor von V1 wird deshalb negativer und damit auch über C1 die Basis von V2, dessen Kollektorstrom zu sinken beginnt. Die deshalb positiver werdende Spannung am Kollektor von V2 wird über C2 auf die Basis von V1 gekoppelt und öffnet diesen vollständig. Über C1 wird nun V2 gesperrt.

Jetzt wiederholt sich der gleiche Vorgang in umgekehrter Reihenfolge. Es entlädt sich C1 über R2 so lange, bis durch V2 ein geringer Strom zu fließen beginnt.

Da beide Stufen fest miteinander gekoppelt sind, erfolgt der Umschaltvorgang in sehr kurzer Zeit. Beide Transistoren „kippen“ aus einem Zustand in den anderen. Die extrem hohe Mitkopplung wird deutlich, wenn man bedenkt, daß nach Verlassen des Übersteuerungsbereichs eine Spannungsänderung an der Basis mit dem Produkt der Spannungsverstärkung beider Transistoren multipliziert als Mitkopplungsspannung erscheint.

Für die Impulsdauer ist die Zeitkonstante aus Koppelkondensator und Basisvorwiderstand bestimmend (*Entladezeitkonstante des Kondensators*). Die Aufladung (auch *Rückladung*) des Kondensators erfolgt, wie Bild 4.29 für C1 zeigt, über die voll geöffnete Basis-Emitter-Diode des Transistors und den Arbeitswiderstand von V1. Die erforderliche Rückladung hat zur Folge, daß die Spannung am Kollektor beim positiven Spannungssprung infolge des durch den Arbeitswiderstand fließenden Ladestroms nicht sofort den Wert der Betriebsspannung erreicht.

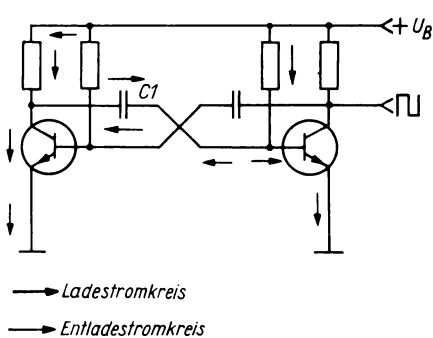


Bild 4.29. Auflade- und Entladestromkreis des Kondensators C1

Damit diese Anstiegsverzögerung im Verhältnis zur Impulsbreite klein wird, muß die Entladezeitkonstante viel größer als die Rückladezeitkonstante sein. Die Basisvorwiderstände dürfen jedoch nur so groß gewählt werden, daß sie einen beim gegebenen Arbeitswiderstand ausreichenden Basisstrom zur völligen Durchsteuerung des Transistors fließen lassen.

Es soll nun anhand des I_C - U_{CE} -Kennlinienfelds des Transistors (hier mit U_{BE} als Parameter) der idealisierte Kurvenverlauf an den einzelnen Schaltungspunkten dargestellt werden (Bild

4.30). Bei der Konstruktion im Kennlinienfeld wird mit dem negativen Spannungssprung an der Basis begonnen. Der Transistor wird damit völlig gesperrt; die Spannung am Kollektor muß ihren Höchstwert erreichen (jedoch durch die erforderliche Rückladung des Kondensators etwas verzögert). Nun wird die Spannung an der betrachteten Basis durch Kondensatorentladung über den Basisvorwiderstand positiver, bis ein geringer Basisstrom zu fließen beginnt. Es erfolgt der bereits beschriebene Umkippvorgang. Dabei wird der Transistor stark übersättigt, was auch zu einer kleinen Kollektorspannungsspitze führt. (Exakt betrachtet haben die Kennlinien im Sättigungsbereich den dargestellten Verlauf; die Sättigungsspannung ist vom Basisstrom abhängig.)

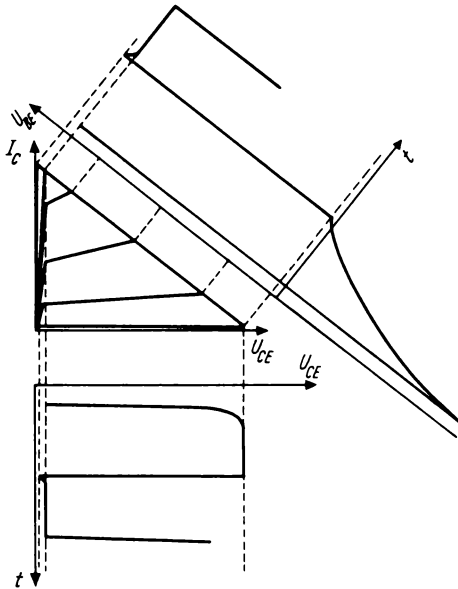


Bild 4.30. Darstellung der Spannungsverläufe einer astabilen Kippstufe im Kennlinienfeld

Im Bild 4.31 sind die Spannungsverläufe an den einzelnen Punkten der Schaltung sowie der Ladespannungsverlauf am Kondensator $C1$ in zeitlicher Abhängigkeit voneinander dargestellt. Zu beachten ist besonders, daß die Sperrspannungsbelastung der Basis-Emitter-Diode fast gleich dem Betrag der Betriebsspannung wird. Sollte die Betriebsspannung größer als diese zulässige Sperrspannung sein, sind zusätzliche Schaltmaßnahmen erforderlich. Eine Variante wird im Bild 4.32 gezeigt.

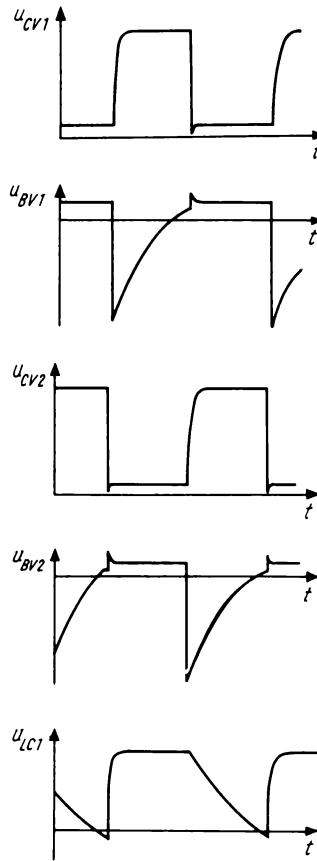


Bild 4.31. Spannungsverläufe an den einzelnen Schaltungspunkten der Schaltung nach Bild 4.28

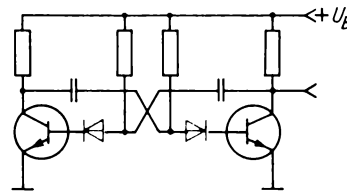


Bild 4.32. Schaltung des astabilen Kippgenerators für höhere Betriebsspannungen

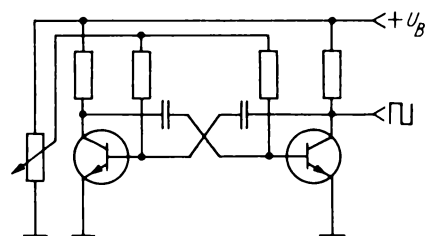


Bild 4.33. Frequenzsteuerbarer astabiler Kippgenerator

Werden die Basisvorwiderstände nicht an die Betriebsspannung, sondern an eine einstellbare Spannung angeschlossen (Bild 4.33), so läßt sich die Frequenz mit dieser Spannung steuern. Die Spannungssprünge an der Basis hängen von den Kollektorspannungssprüngen ab und werden nicht von der an die Basisvorwiderstände angelegten Spannung beeinflusst. Hingegen hängt die Entladezeit von dieser Spannung ab und läßt sich somit steuern.

4.3.3. Sägezahngeneratoren

In vielen Geräten werden Spannungen mit zeitlinearem Anstieg benötigt (z. B. Horizontalablenkspannung in Oszilloskopen und Fernsehempfängern, Analog-Digital-Wandler). Eine ansteigende Spannung tritt während des Ladevorgangs am Kondensator auf (s. Bild 4.22). Der Spannungsanstieg erfolgt hierbei allerdings nach einer Exponentialfunktion. Um diese Spannung sich rhythmisch wiederholend gewinnen zu können, muß der Kondensator nach der Ladung in möglichst kurzer Zeit wieder entladen werden, um danach erneut aufgeladen werden zu können.

Für die Linearisierung des exponentiellen Spannungsanstiegs gibt es verschiedene Möglichkeiten. So ist es möglich, dafür zu sorgen, daß nur ein sehr kleiner Teil der Ladekurve durchlaufen wird, der dann als weitgehend linear betrachtet werden kann. Dieses Prinzip wird bei einem an eine Rechteckspannung angeschlossenen entsprechend dimensionierten Tiefpaß ausgenutzt, der, dann als *Integrierglied* bezeichnet, bereits im Zusammenhang mit Bild 4.25 erläutert wurde. Diese Möglichkeit hat jedoch den Nachteil, daß die Amplitude der Ausgangssägezahnspannung sehr viel kleiner als die Amplitude der angelegten Rechteckspannung wird.

Eine weitere Möglichkeit der Linearisierung des Spannungsanstiegs ergibt sich durch die Beziehungen

$$C = \frac{\Delta Q}{\Delta U_C} \text{ und } I = \frac{\Delta Q}{\Delta t}.$$

Daraus folgt

$$\Delta U_C = \frac{1}{C} \Delta Q$$

$$\Delta U_C = \frac{1}{C} I \Delta t. \quad (4.9)$$

Man erkennt, daß bei konstantem Strom I die Spannung am Kondensator linear mit der Zeit ansteigt. Zur Konstanthaltung des Stromes I während der Aufladung des Kondensators gibt es mehrere Varianten.

Zuerst soll der mit dem Kondensator in Reihe geschaltete Widerstand durch eine *Konstantstromquelle* ersetzt werden. Eine entsprechende Schaltung wird im Bild 4.34 gezeigt. Der als Konstantstromquelle geschaltete Transistor $V1$ erhält über den Basisspannungsteiler ein festes Basispotential. Der Emittierwiderstand bewirkt die Gegenkopplung und damit den konstanten Strom durch $V1$. Durch den als Schalter arbeitenden Transistor $V2$ wird der Kondensator nach der Aufladung kurzzeitig entladen. Während der Sperrzeit von $V2$ erfolgt die zeitlineare Aufladung des Kondensators.

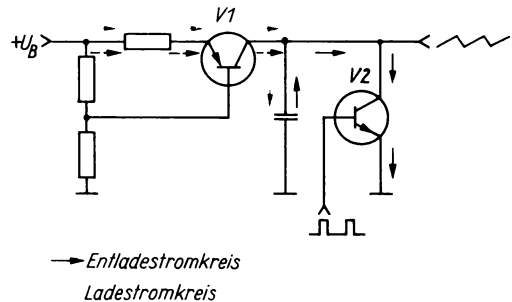


Bild 4.34. Linearisierung des Spannungsanstiegs durch Konstantstromquelle

Der durch den Widerstand (s. Bild 4.21) fließende Strom ergibt sich zu

$$i_C = \frac{U - u_C}{R}. \quad (4.10)$$

Sorgt man nun dafür, daß die Differenz $U - u_C$ während der Ladezeit durch eine Erhöhung der Spannung U um u_C konstant bleibt, so wird auch der Ladestrom konstant bleiben (mitlaufende Ladespannung). Die mitlaufende Ladespannung zur Erzeugung einer Spannung mit zeitlinearem Anstieg wird z. B. in der Bootstrapp-Schaltung genutzt (Bild 4.35). Das verstärkende Bauelement muß hierin eine Verstärkung von 1 haben, und der in die Eingangsklemme hineinfließende Strom i_1 sollte vernachlässigbar klein sein gegenüber i_C . Diese Bedingung erfüllt ein Spannungsfolger oder eine Kollektorschaltung. Unter der gegebenen Voraussetzung folgt für i

$$i = \frac{u_R}{R}$$

und mit u_R aus dem Maschensatz (Maschenumlauf im Bild mit Strichlinie dargestellt)

$$u_R = U + u_2 - u_1 = U + u_1 (v_u - 1)$$

$$i = \frac{U}{R} + \frac{(v_u - 1) u_1}{R} \quad (4.11)$$

Mit $v_u = 1$ folgt weiter

$$i = \frac{U}{R}$$

Der Strom bleibt während der Ladezeit konstant. Die Entladung erfolgt durch Schließen des Schalters.

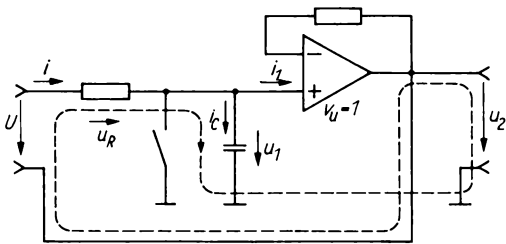


Bild 4.35. Prinzip der Bootstrap-Schaltung

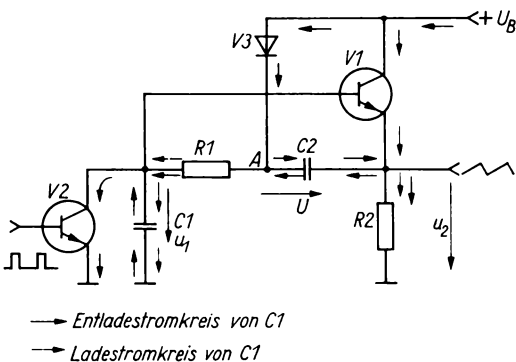


Bild 4.36. Bootstrap-Schaltung mit Transistoren

Eine mit Transistoren realisierte Schaltung wird im Bild 4.36 gezeigt. Der Transistor $V2$ übernimmt die Funktion des gesteuerten Schalters zur Entladung von $C1$; der $V1$ arbeitet in Kollektorschaltung ($v_u \approx 1$, Eingangswiderstand sehr hoch); $C2$ ergibt mit seiner Ladespannung die Spannung U im Bild 4.35. Zur Vereinfachung sei zunächst angenommen, daß $C2$ auf die konstante Spannung U aufgeladen sei. Für i gilt dann wieder Gl. (4.11), und mit $v_u \approx 1$ bleibt i konstant. Zur Erklärung der Einzelvorgänge nehmen wir an, daß $C1$ gerade über $V2$ entladen werde. Dann ist auch $V1$ gesperrt, und

es wird $C2$ über $V3$ und $R2$ etwa auf den Wert der Betriebsspannung U_B aufgeladen. Öffnet jetzt $V2$, so wird $C1$ durch die Ladespannung des Kondensators $C2$ aufgeladen und $V1$ geöffnet. Zur Ladespannung von $C2$ addiert sich nun u_2 (mitlaufende Ladespannung). $V3$ bleibt während dieser Zeit gesperrt, da durch die Addition von u_2 zur Ladespannung U Punkt A positiver als $+U_B$ wird. Durch erneutes Öffnen von $V2$ wird der Ladevorgang beendet, und $C1$ wird wieder entladen. Gleichzeitig wird dem Kondensator $C2$ über $V3$ die verlorengegangene Ladung wieder zugeführt. Damit die Ladespannung von $C2$ während des Ladevorgangs von $C1$ sich möglichst nur wenig ändert, muß die Kapazität von $C2$ wesentlich größer als die von $C1$ gewählt werden. Außerdem muß $R2$ sehr viel kleiner als $R1$ sein. Im Bild 4.36 sind die während der Lade- und Entladezeit von $C1$ fließenden Ströme mit Strompfeilen dargestellt.

Das Prinzip der mitlaufenden Ladespannung wird auch beim Miller-Integrator angewendet (Bild 4.37). Beim Miller-Integrator ist ein invertierender Verstärker mit möglichst hoher Verstärkung und möglichst hohem Eingangswiderstand erforderlich. Die vorgegebenen Richtungen von u_1 und u_2 berücksichtigen die Phasendrehung des Verstärkers.

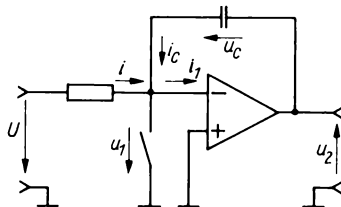


Bild 4.37. Miller-Integrator

Für die Ströme gilt: $i_1 \ll i$ und damit für den Knotenpunkt $i \approx -i_c$.

Für die Spannungen gilt: u_1 wird vernachlässigbar klein und damit $u_2 \approx -u_c$.

Um den fast zeitlinearen Anstieg der Spannung u_2 (ungefähr gleich u_c) zu begründen, soll der Strom i am Anfang und am Ende der Ladung berechnet werden. Vor der Ladung von C gilt (Schalter geschlossen)

$$i = \frac{u_R}{R} = \frac{U - u_1}{R}$$

und mit $u_1 = 0$

$$i = \frac{U}{R} \quad (4.12)$$

Am Ende der Ladung gilt (Schalter geöffnet)

$$i = \frac{U - u_1}{R}. \quad (4.13)$$

Dabei ist nun u_1 in Abhängigkeit von u_2 und v_u zu ermitteln. Unter Verwendung des Maschensatzes gilt

$$u_1 = -u_2 - u_C; \quad u_2 = v_u u_1$$

$$u_1 = -v_u u_1 - u_C$$

$$u_1 = -\frac{u_C}{1 + v_u}.$$

Da wegen der hohen Spannungsverstärkung auch gilt

$$u_C \approx -u_2,$$

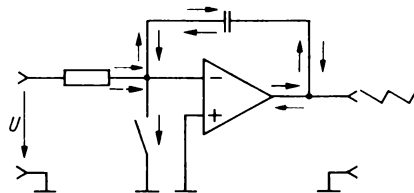
folgt für u_1 :

$$u_1 = \frac{u_2}{1 + v_u}.$$

Dieser Ausdruck wird in Gl. (4.13) eingesetzt, und es folgt für den Strom i am Ende der Ladung:

$$i = \frac{U}{R} - \frac{u_2}{(1 + v_u) R}. \quad (4.14)$$

Bei genügend hoher Verstärkung ist der Strom am Ende der Ladung nur geringfügig kleiner als am Anfang der Ladung. Bild 4.38 zeigt die fließenden Ströme während der Auf- und Entladung des Kondensators.



→ Entladestromkreis
→ Ladestromkreis

Bild 4.38. Ströme im Miller-Integrator

Eine weitere wichtige Eigenschaft des Miller-Integrators ist die Vergrößerung der Zeitkonstante bestimmten Kapazität. Der Kondensator C erhält in der Schaltung eine Ladung entsprechend seiner Spannung u_C . Für diese Spannung gilt nach Bild 4.37

$$u_C = -u_2 - u_1; \quad u_2 = v_u u_1$$

$$u_C = -u_1 (1 + v_u).$$

Dem Kondensator wird Ladung mit einer mit $(1 + v_u)$ multiplizierten Spannung u_1 zugeführt.

Denkt man sich den Kondensator gemäß Bild 4.39 als C_{ers} zwischen Eingang und Masse geschaltet, muß dieser zur Aufnahme der gleichen Ladung folgende Kapazität haben:

$$|Q| = C_{ers} u_1$$

$$|Q| = C u_1 (1 + v_u)$$

$$C_{ers} = C (1 + v_u). \quad (4.15)$$

Im Miller-Integrator erfolgt, auf den Verstärkereingang bezogen, eine Vergrößerung der Kapazität C um den Faktor $(1 + v_u)$.

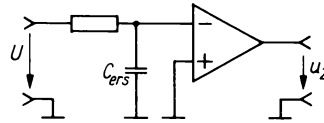
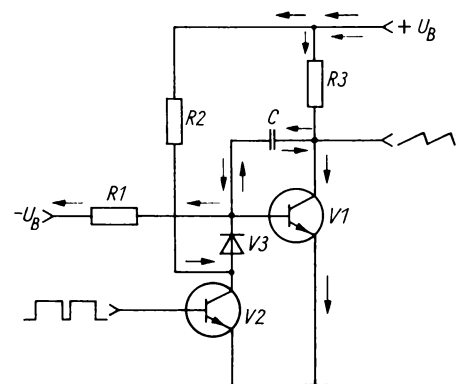


Bild 4.39. Darstellung der Ersatzkapazität im Miller-Integrator

Ein mit Transistoren realisierter Miller-Integrator ist im Bild 4.40 dargestellt. Darin übernimmt $V1$ die Funktion des invertierenden Verstärkers, $V2$ mit $V3$ die Schalterfunktion. Ist $V2$ gesperrt, wird $V1$ über $R2$ und $V3$ voll geöffnet. Am Ausgang der Schaltung liegt die Sättigungsspannung von $V1$, und der Kondensator ist weitgehend entladen. Wird jetzt $V2$ leitend, so hat er durch die dann gesperrte Diode $V3$ keinen Einfluß mehr auf die Schaltung. Es erfolgt die Aufladung des Kondensators. Dabei wird die Spannung an der Basis von $V1$ langsam negativer, und die Ausgangsspannung steigt linear an. Der Vorgang wird durch erneutes Sperren von $V2$ und die damit verbundene Entladung von C beendet. Die Stromverläufe während des La-



→ Entladestromkreis
→ Ladestromkreis

Bild 4.40. Miller-Integrator mit Transistoren

dens und Entladens von C sind im Bild 4.40 eingetragen.

Ein Beispiel eines freilaufenden Sägezahngenerators ist die *Sperrschwingerschaltung* (Bild 4.41). Der Basisspannungsteiler ist dabei so dimensioniert, daß der Transistor nach erfolgter Aufladung von $C1$ voll durchsteuert wäre. Im Einschaltmoment ist $C1$ entladen; damit ist der Transistor gesperrt. Steigt jetzt langsam die Ladespannung an $C1$, so wird nach Erreichen der Schwellspannung des Transistors ein Kollektorstrom zu fließen beginnen, der in der Mitkopplungsspule eine Spannung induziert, die infolge eines sehr hohen Mitkopplungsgrads eine völlige Durchsteuerung des Transistors bewirkt. (Die Mitkopplungsspannung addiert sich zur Ladespannung von $C1$.) Ist der Kollektorstrom auf seinen Höchstwert gestiegen und kein weiterer Anstieg mehr möglich, so wird auch die Mitkopplungsspannung wieder Null. Die Schaltung schwingt in den Sperrzustand zurück. Dabei wirkt die jetzt geöffnete Diode verzögernd auf den Abbau der magnetischen Feldenergie und verhindert die Ausbildung unzulässig hoher Spannungsspitzen. Gleichzeitig wird der Kondensator $C1$ durch die während des Stromanstiegs in der Mitkopplungsspule induzierte Spannung so stark negativ aufgeladen, daß ein erneuter Stromanstieg zunächst nicht möglich ist. Es muß jetzt erst der Kondensator $C1$ so lange umgeladen werden, bis an der Basis des Transistors wieder die Schwellspannung erreicht ist. Im Bild 4.42 sind die Spannungsverläufe an der Basis (u_B), am Kollektor (u_C) und am Kondensator (u_{C1}) sowie der Kollektorstromverlauf in ihrer zeitlichen Zugehörigkeit vereinfacht dargestellt. Der Auf- und Entladestromkreis von $C1$ ist im Bild 4.41 zusätzlich eingetragen.

Zusammenfassung

Die Art der Verformung von rechteckförmigen Spannungen an RC -Schaltungen hängt vom Größenverhältnis zwischen der Zeitkonstanten τ der RC -Schaltung und der Periodendauer T der Rechteckspannung ab.

Der astabile Kippgenerator ist ein 2stufiger RC -Verstärker mit direkter Kopplung zwischen Ausgang und Eingang. Die Anstiegsflanke der Rechteckspannung wird durch die Rückladezeitkonstante bestimmt. An der Basis tritt eine Sperrspannungsbelastung auf, die etwa gleich der Betriebsspannung ist. Eine Steuerung der Frequenz ist möglich, wenn man die Basisvor-

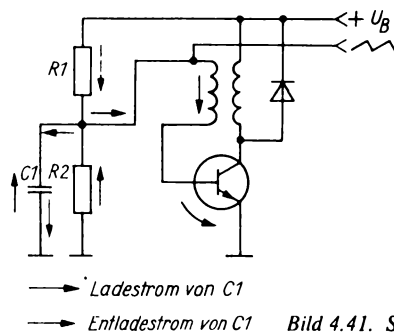


Bild 4.41. Sperrschwinger

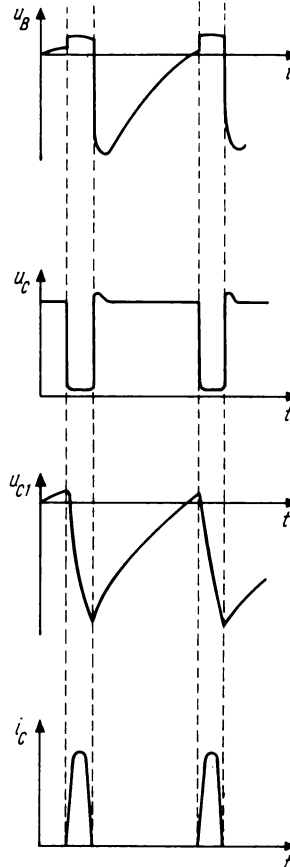


Bild 4.42. Strom- und Spannungsverläufe am Sperrschwinger

widerstände an eine einstellbare Spannung anschließt.

Eine Spannung mit zeitlinearem Anstieg (Sägezahnspannung) wird aus dem Aufladespannungsverlauf eines Kondensators abgeleitet. Entweder nutzt man nur einen sehr kleinen Teil der Ladekurve, oder man sorgt für eine Konstanthaltung des Ladestroms, um eine genügende Linearität des Spannungsanstiegs zu er-

halten. Einen konstanten Ladestrom erreicht man entweder durch Einsatz einer Konstantstromquelle oder durch die mitlaufende Ladespannung. Diese wird in der Bootstrap-Schaltung und im Miller-Integrator angewendet.

4.4.

Aufgaben

1. Leiten Sie die Mitkopplungsbedingung für einen rückgekoppelten Verstärker ab!
2. Erklären Sie die Notwendigkeit der Amplitudenbegrenzung!
3. Wie wirkt die Sättigungsbegrenzung?
4. Konstruieren Sie punktweise die dynamische Stromverstärkungskennlinie eines Transistors aus dem Ausgangskennlinienfeld!
5. Unter welchen Bedingungen entsteht eine Sinusspannung?
6. Warum entsteht für die Bedingung $k v_u \gg 1$ keine Sinusspannung?
7. Wie wird beim Meißner-Oszillator die erforderliche Phasendrehung realisiert?
8. Wodurch unterscheiden sich Serien- und Parallelspeisung?
9. Durch welche Maßnahmen kann der Klirrfaktor bei Transistoroszillatoren erniedrigt werden?
10. Warum entsteht beim Meißner-Oszillator eine sinusförmige Spannung?
11. Erklären Sie die Entstehung der Phasendrehung im Mitkopplungsvierpol bei der induktiven Dreipunktschaltung!
12. Begründen Sie, warum die Betriebsspannung wechsellspannungsmäßig Massepotential hat!
13. Warum bevorzugt man für hohe Frequenzen die Basisschaltung, und warum muß dabei die Basis wechsellspannungsmäßig Massepotential haben?
14. Erklären Sie die vorhandene Widerstandsanpassung zwischen Ausgangs- und Eingangswiderstand anhand der Schaltung nach Bild 4.11!
15. Warum ist mit Quarzoszillatoren eine hohe Frequenzkonstanz erreichbar?
16. Welche Aufgabe haben zusätzliche Schwingkreise in Quarzoszillatoren?
17. Welche Schaltungen werden mit den Begriffen VCO und VCXO bezeichnet, und welche Eigenschaften haben diese Schaltungen?
18. Warum sind beim RC-Oszillator mindestens drei RC-Glieder erforderlich?
Hinweis: Führen Sie den Nachweis grafisch durch!
19. Welche Nachteile entstehen bei einer zu großen Anzahl von RC-Gliedern?
20. Warum ist die Selbsterregungsbedingung beim RC-Oszillator mit Phasenschieberkette nur für eine Frequenz erfüllt?
21. Welche Abhängigkeit besteht zwischen dem Mitkopplungsfaktor k und der Anzahl der RC-Glieder?
(Durch Konstruktion der Zeigerbilder ist jeweils die angenäherte Größe von k zu bestimmen!)
22. Wie unterscheidet sich das reale Zeigerbild der Phasenschieberkette vom vereinfachten Zeigerbild?
23. Skizzieren Sie das Zeigerbild des Wien-Brücken-Oszillators sowohl für eine größere als auch für eine kleinere Frequenz gegenüber der für das Zeigerbild Bild 4.17 b gewählten Frequenz!
24. Erläutern Sie die Wirkungsweise der Schaltung nach Bild 4.19!
25. Was versteht man unter der Zeitkonstanten einer RC-Schaltung?
26. Wie wird τ berechnet?
27. Kontrollieren Sie, daß für $\tau \ll t_i$ die Summe der für beide Schaltungen (Bilder 4.25 und 4.27) gezeigten Ausgangsspannungsverläufe gleich der Eingangsspannung ist!
28. Führen Sie dieselbe Betrachtung für $\tau \gg t_i$ durch!
29. Bestimmen Sie für einen Kondensator von 10 nF und eine Frequenz der Eingangsspannung von 1 kHz die erforderlichen Widerstände für $\tau \ll t_i$; $\tau \approx t_i$; $\tau \gg t_i$!
30. Warum sind am Kondensator keine sprunghaften Spannungsänderungen möglich?
31. Erklären Sie die Wirkungsweise des astabilen Kippgenerators!
32. Orientieren Sie sich anhand der Fachliteratur über weitere Schaltungsvarianten, die die Nachteile der hier genannten Schaltung (Bild 4.28) nicht haben!
33. Interpretieren Sie die Spannungsverläufe an den einzelnen Punkten der Schaltung nach Bild 4.28!
34. Welche Möglichkeiten zur Erzeugung einer zeitlinear ansteigenden Spannung gibt es?
35. Begründen Sie, wieso in der Schaltung nach Bild 4.34 der Ladestrom konstant gehalten wird!
36. Warum erfolgt der Spannungsanstieg am Kondensator der Bootstrap-Schaltung zeitlinear?
37. Wie wird in der Schaltung nach Bild 4.36 erreicht, daß der Kondensator C_I in gegenüber der Aufladezeit kurzer Zeit entladen wird?
38. Warum ist die Spannung u_2 im Bild 4.36 fast gleich der Spannung u_1 ?
39. Erklären Sie die Wirkungsweise des Miller-Integrators!
40. Wie entsteht beim Miller-Integrator die mitlaufende Ladespannung?
41. Warum entsteht am Ausgang der Schaltung nach Bild 4.38 eine zeitlinear abfallende Spannung?
42. Erklären Sie die Wirkungsweise eines Sperrschwingers!
43. Wodurch unterscheidet sich der Sperrschwinger von einem Meißner-Oszillator?

5.

Elektronische Grundsaltungen der Digitaltechnik

5.1.

Allgemeines

In den letzten Jahren ist die Elektronik in alle Zweige der Wirtschaft und Wissenschaft eingedrungen. Neben den bisher bekannten Einsatzgebieten haben dabei elektronische Schaltungen zur Datenverarbeitung und zur Automatisierung von Produktionsprozessen besondere Bedeutung erlangt. Für die hier gestellten Aufgaben sind digitale elektronische Schaltungen gut geeignet. Sie sind aus Grundsaltungen zusammengesetzt, bei denen nur zwei mögliche Ausgangssignale zu unterscheiden sind. Die Bezeichnung dieser Ausgangssignale erfolgt entweder als *Spannungspegel* mit den Buchstaben H (high, hoch) und L oder T (low, tief) oder als *logischer Wert* bzw. *logischer Zustand* mit den Bezeichnungen 0 und 1. Eine Zuordnung der einzelnen Schreibweisen wird in Tafel 5.1 gezeigt. Im Bild 5.1 sind Beispiele von Signalpegeln mit den zugehörigen logischen Werten eingetragen. Man kann feststellen, daß jeweils ein großer

Spannungsbereich durch die Schaltung eindeutig dem Pegel H oder L zugeordnet wird. Der Einfluß von Bauelementetoleranzen ist damit wesentlich geringer als bei analogen Schaltungen. Die zulässigen größeren Toleranzen und die Möglichkeit, selbst komplizierte digitale Schaltungen (z. B. Mikroprozessoren) auf wenige Grundsaltungen zurückführen zu können, bedeuten, daß digitale Schaltungen integrationsfreundlich sind. So ist auch zu erklären, daß die Entwicklung und Fertigung von integrierten Schaltungen in großem Umfang zuerst bei digitalen Schaltungen begann. Da innerhalb digitaler Schaltungen nur zwei Pegel unterschieden werden, lassen sich auch Amplitudenstörungen, die auf dem Übertragungsweg entstanden sind, verhältnismäßig einfach beseitigen. Das führte in den letzten Jahren verstärkt zum Einsatz digitaler Schaltungstechnik für die Übertragung analoger Signale. So wird beispielsweise bei der PCM (Pulsmodulation) ein zu übertragendes Sprachsignal digitalisiert und in Form einer Impulsfolge übertragen. Auch in der NF-Technik sind neue Verfahren entwickelt worden, bei denen ein NF-Signal beispielsweise digitalisiert auf Platte oder Band gespeichert wird. Damit ist eine wesentliche Rauschminderung erzielbar. Darüber hinaus werden in der Meßtechnik heute Meßwerte vielfach digital mit hoher Genauigkeit erfaßt. Ebenfalls völlig neue Möglichkeiten eröffnet der Einsatz der Digitaltechnik bei der Ableitung belie-

Tafel 5.1. Bezeichnung binärer Signale

Pegel	Logischer Wert	
	positive Logik	negative Logik
H (positiverer Pegel)	1	0
L oder T (negativerer Pegel)	0	1

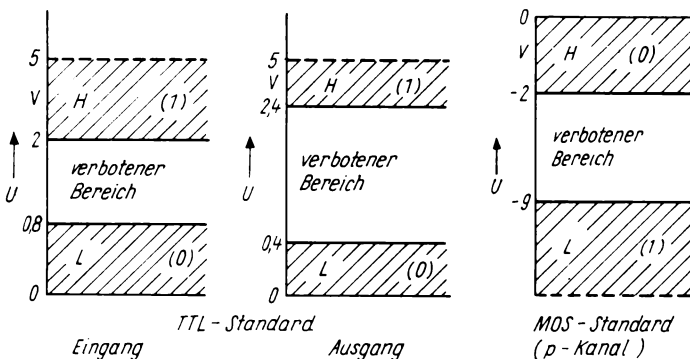


Bild 5.1. Signalpegel

biger quartzgenauer Frequenzen von nur einer festliegenden quartzgenauen Frequenz. Dieses Verfahren – die Frequenzsynthese – wird zunehmend zur Frequenzeinstellung bei Sendern und Empfängern eingesetzt.

Gegenüber mit Kontakten ausgeführten Schaltungen zeichnen sich elektronische durch die höhere Schaltgeschwindigkeit und eine praktisch unbegrenzte Schalthäufigkeit aus. Außerdem ist es durch die Entwicklung integrierter Schaltkreise bei hoher Zuverlässigkeit möglich geworden, selbst komplizierte Schaltungen preiswert mit niedriger Verlustleistung und niedrigem Platzbedarf zu fertigen (z. B. Taschenrechner).

Die verwendeten Schaltungen setzen sich aus für das jeweilige System typischen Grundgattern zusammen, von denen einige nachfolgend erläutert werden sollen. Die innere Schaltung von Schaltkreisen mit höherem Integrationsgrad ist für den Anwender oft kaum übersehbar – ihn interessieren das in verschiedenen Formen beschriebene Verhalten der IS und die Zusammenschaltungsbedingungen.

5.2.

Transistor als elektronischer Schalter

Der Transistor kann als *elektronischer Schalter* eingesetzt werden. Bild 5.2 zeigt die dabei auftretenden Arbeitspunkte im Kennlinienfeld. Liegt der Arbeitspunkt A vor, so wird der Transistor von einem hohen Strom durchflossen, und es fällt an ihm nur eine sehr geringe Spannung ab. Der Transistor wirkt damit wie ein ge-

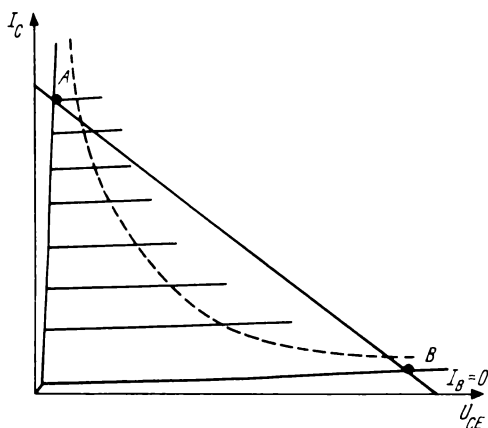


Bild 5.2. Darstellung der Arbeitspunkte einer Transistor-schaltstufe

schlossener Schalter. Im Arbeitspunkt B fällt die volle Betriebsspannung am Transistor ab; durch ihn fließt nur der sehr geringe Reststrom. Der Transistor wirkt damit wie ein geöffneter Schalter. Erfolgt der Übergang von Arbeitspunkt A nach Arbeitspunkt B in sehr kurzer Zeit, so darf die Widerstandsgerade die Verlustleistungshyperbel schneiden. Der Arbeitspunkt A wird dann durch den maximal zulässigen Kollektorstrom und der Arbeitspunkt B durch die maximal zulässige Kollektor-Emitter-Spannung begrenzt.

Beispiel

Der Transistor SF 827 hat folgende Grenzwerte:

$U_{CE0} = 30 \text{ V}$; $I_C = 0,5 \text{ A}$; $P_{tot} = 735 \text{ mW}$ bei $t_a = 25^\circ \text{C}$.

Maximal ist mit ihm ein Verbraucher mit 30 V Betriebsspannung und 0,5 A Stromaufnahme schaltbar. Das entspricht einer Schaltleistung von 15 W.

In der Praxis kann dieser Wert nur selten voll genutzt werden. Das Beispiel zeigt, daß ein Transistor eine gegenüber seiner Verlustleistung wesentlich höhere Leistung schalten kann.

Um den Arbeitspunkt A sicher zu erreichen, arbeitet man mit einer etwa 1,5fachen Übersteuerung. Würde nämlich der Basisstrom nur geringfügig zu niedrig werden, so müßte der Arbeitspunkt in den durch die Verlustleistungshyperbel abgetrennten Bereich zu hoher Verlustleistung gelangen. Die Folge wäre eine Zerstörung des Transistors.

(Bei der Bestimmung des erforderlichen Basisstromes ist zu beachten, daß h_{21E} stark vom Kollektorstrom abhängt und die für den Transistor angegebene Stromverstärkung folglich nur bei den angegebenen Meßbedingungen Gültigkeit hat. Ist beispielsweise der bei den Meßbedingungen angegebene Strom kleiner als der zu schaltende Strom, ist mit einer Verringerung der Stromverstärkung gegenüber den angegebenen Werten zu rechnen.)

Um im Arbeitspunkt B einen möglichst geringen Stromfluß durch den Transistor zu erhalten, kann man bei Leistungsschaltstufen die Basis-Emitter-Strecke in Sperrichtung vorspannen. Diese Vorspannung in Sperrichtung ist nur unter Zuhilfenahme einer zweiten Betriebsspannungsquelle mit entgegengesetzter Polarität möglich.

Oft kommt es vor, daß *Signallampen* mit dem Transistor zu schalten sind. Die hierbei auftretenden *Einschaltströme* (bedingt durch den gegenüber dem Betriebswiderstand sehr niedrigen

Kaltwiderstand der Glühlampe) dürfen den höchstzulässigen Kollektorstrom des Transistors nicht überschreiten. Um trotzdem mit Transistoren kleiner Verlustleistung eine hohe Glühlampenleistung schalten zu können, muß man entweder die Einschaltstromspitzen durch einen Vorwiderstand begrenzen (Bild 5.3a) oder die Glühlampe nach Bild 5.3b vorheizen. Der Widerstand wird dabei so festgelegt, daß bei gesperrtem Transistor der Glühfaden fast unsichtbar glüht. Wird ein *Relais* oder *Schütz* durch den Transistor geschaltet, so ist besonders die hohe *Selbstinduktionsspannungsspitze* beim Abschalten zu beachten. Hier schaltet man eine Diode parallel zur induktiven Last (Bild 5.4). Die Wirkungsweise dieser Schaltung ist im Abschnitt 2.2. beschrieben.

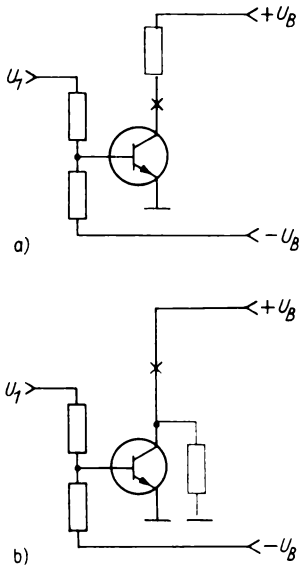


Bild 5.3. Begrenzung des Einschaltstromstoßes von Glühlampen

a) mit Vorwiderstand; b) mit Vorheizwiderstand

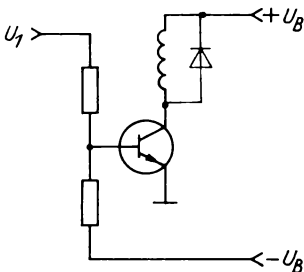


Bild 5.4. Schaltstufe mit induktiver Last

Zusammenfassung

Ein Transistor kann als elektronischer Schalter benutzt werden. Wenn der Übergang vom Ein- zum Auszustand in sehr kurzer Zeit erfolgt, kann eine wesentlich über der maximalen Verlustleistung liegende Leistung geschaltet werden. Der Eingangsspannungsteiler setzt die durch den Signalpegel bestimmte Eingangsspannung in die erforderliche Basis-Emitter-Spannung um. Bei induktiver Last ist die Selbstinduktionsspannungsspitze beim Abschalten und bei Signallampen als Last die Einschaltstromspitze zu beachten und wirkungsvoll zu begrenzen.

5.3.

Kombinatorische Schaltungen

5.3.1.

Allgemeines

Bei kombinatorischen Schaltungen ist die Ausgangsgröße nur vom augenblicklichen Zustand der Eingangsgrößen abhängig. Rückführungen von Ausgängen auf davorliegende Eingänge sind nicht möglich. Das Verhalten der logischen Elemente wird durch die Wahrheitstafel (auch Wertetafel oder Funktionsmatrix genannt) oder eine Gleichung beschrieben.

Betrachten wir das Zusammenwirken von logischen Elementen an einem einfachen Beispiel. Die Sicherungsanlage eines Reaktionsgefäßes besteht aus drei gleichwertigen Gefahranzeigegeräten. Der Reaktionsvorgang soll abgeschaltet werden, wenn mindestens zwei der drei Gefahrensignale anliegen. Die Gleichung dieser Schaltung lautet:

$$y = x_1 x_2 \vee x_1 x_3 \vee x_2 x_3.$$

Bild 5.5 zeigt die zugehörige Schaltung.

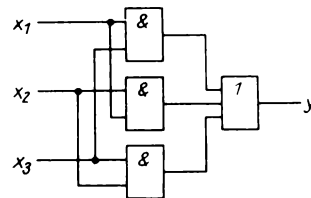
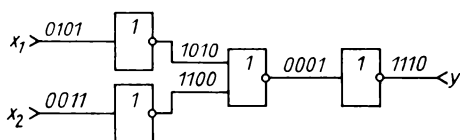


Bild 5.5. Zusammenschaltungsbeispiel

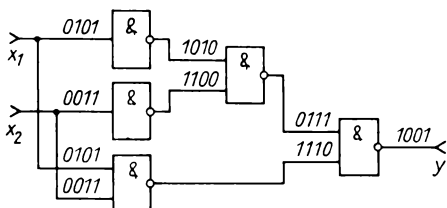
Im Bild 5.6 ist ein weiteres einfaches Beispiel einer kombinatorischen Verknüpfung dargestellt. Um die Wahrheitstabelle dieser Schaltung ermitteln zu können, trägt man sich an

den Eingängen alle möglichen Kombinationen logischer Werte ein und danach die daraus folgenden logischen Werte an allen anderen Punkten der Schaltung. Im Bild 5.6 entspricht die Reihenfolge der logischen Werte von links nach rechts den Zeilen von oben nach unten in der Wahrheitstabelle. Am Beispiel ist zu erkennen, daß eine NAND-Funktion durch ausschließlich NOR-Elemente realisierbar ist. Selbstverständlich sind auch andere Grundsaltungen ausschließlich durch NAND oder durch NOR realisierbar. So zeigt Bild 5.7 die Realisierung der Äquivalenz durch ausschließlich NAND-Elemente.



x_1	x_2	y
0	0	1
1	0	1
0	1	1
1	1	0

Bild 5.6. Realisierung einer NAND-Funktion durch ausschließlich NOR-Elemente



x_1	x_2	y
0	0	1
1	0	0
0	1	0
1	1	1

Bild 5.7. Realisierung einer Äquivalenzfunktion durch ausschließlich NAND-Elemente

5.3.2.

Grundgatter

Die logischen Elemente lassen sich schaltungs-technisch in vielfältigen Formen realisieren. Von besonderer Bedeutung sind Varianten in integrierter Technik, die in den folgenden Abschnitten ausschließlich behandelt werden. Neben ökonomischen Gesichtspunkten sind für

die Auswahl in der Praxis technische Gesichtspunkte zu beachten, wie Schaltzeit, Verlustleistung, Störsicherheit und Anpaßbarkeit an die Anlage. Die rasche Weiterentwicklung integrierter digitaler Schaltungen führt zu Schwerpunktverlagerungen des Einsatzes der einzelnen Techniken.

Tafel 5.2 enthält eine Zusammenstellung wichtiger technischer Parameter einiger Schaltkreisfamilien. Deutlich ist innerhalb der Gruppe der TTL-Schaltungen zu erkennen, daß eine Verkürzung der Schaltzeit mit einer Verlustleistungserhöhung verbunden ist. Erst durch Änderung der Schaltungsstruktur gelang es mit der TTL-LS, bei gleichbleibender Schaltzeit die Verlustleistung zu erniedrigen. Betrachtet man die Verlustleistung je Gatter, so fällt besonders auf, daß die Verlustleistung der CMOS um drei Zehnerpotenzen und die der P^L um bis zu sechs Zehnerpotenzen niedriger ist als die der TTL-Standards.

TTL-Gatter

Die TTL-Technik baut in ihrer Grundsaltung auf einem Mehremittertransistor auf, der in Verbindung mit dem gesamten Schaltkreis ein NAND-Gatter ergibt. Beim Eingangsmehremit-

Tafel 5.2. Ausgewählte technische Parameter wichtiger Schaltkreisfamilien

Bezeichnung	Bauelementenfunktionen je Grundgatter	Schaltzeit in ns	Verlustleistung je Grundgatter in mW
TTL-Standard	9	10	10
TTL-L low-power-TTL langsame TTL	9	33	1
TTL-H high-speed-TTL schnelle TTL	12	6	23
TTL-LS low-power-Schottky-TTL	16	9	2
CMOS-Standard	4 (8)	20...100	$(1...20) \cdot 10^{-3}$ 0,7 bei 1 MHz
CMOS-HCT	4 (8)	< 10	$> 25 \cdot 10^{-3}$ 1 bei 1 MHz
ECL	6	0,8...2	25...60
I ² L	2	> 10	$> 10 \cdot 10^{-6}$

tertransistor sind alle Emmitter elektrisch gleichwertig und können jeweils für sich allein oder gemeinsam die Emmitterfunktion übernehmen. Die TTL-Technik ist eine *Übersteuerungstechnik*: Um die Funktion der Grundsaltung erklären zu können, bezieht man sich auf charakteristische Werte am übersteuerten Siliziumtransistor. Es beträgt hierbei $U_{BE} \approx 0,7\text{ V}$ und $U_{CE\text{ sat}} \approx 0,2\text{ V}$; $U_{CE\text{ sat}}$ ist die Kollektor-Emittersättigungsspannung, die vom Hersteller für einen bestimmten Arbeitspunkt angegeben wird und beispielsweise dem Arbeitspunkt A im Bild 5.2 entsprechen könnte. Bei inversem Betrieb des Transistors (offener Emmitter, Basis positiver als Kollektor) beträgt U_{BC} ebenfalls etwa $0,7\text{ V}$. Der Kollektorstrom wird dabei negativ. (Er fließt aus dem Kollektor heraus.)

Im Bild 5.8 ist die komplette Schaltung eines TTL-NAND-Gatters mit *Gegentaktausgangsstufe* dargestellt (entsprechend dem IS D100). Sind ein oder mehrere Emmitter auf Pegel L geschaltet, so ergeben sich die möglichen Stromflüsse und Potentiale nach Bild 5.8a. Durch $V1$ fließt ein nur durch $R1$ begrenzter Basisstrom. An der Basis stellt sich ein Potential von etwa $0,7\text{ V}$ ein. Ein Kollektorstrom kann trotzdem nicht fließen, da es vom Kollektor keine leitende Verbindung zu $+U_B$ gibt. Damit steht auch an der Basis von $V2$ das für dessen Öffnung erforderliche Potential von $+0,7\text{ V}$ nicht zur Verfügung, und $V2$ ist gesperrt. Über $R2$ wird $V3$ geöffnet, und es kann je nach angeschlossener Last ein Strom fließen (Emitterstrom von $V3$). Dabei ist das Potential am Ausgang mindestens um die Flußspannung der Basis-Emitter-Diode von $V3$ und der in Reihe liegenden Diode $V5$ niedriger als die Betriebsspannung. Am Ausgang liegt Pegel H. Es wirkt ein *niedriger Ausgangswiderstand* ($V3$ voll durchsteuert). Besonders bei kapazitiver Last können so die Kondensatoren schnell umgeladen werden.

Im Bild 5.8b sind alle Eingänge auf Pegel H geschaltet (hier auf $+U_B$). Durch $V1$ kann kein Emmitterstrom fließen. $V1$ wird invers betrieben, und es fließt der eingezeichnete Strom, der $V2$ öffnet. Dessen Emmitterstrom durchfließt $R3$ und öffnet mit dem hier auftretenden Spannungsabfall $V4$. Der Spannungsabfall an $R2$ bewirkt eine Sperrung von $V3$. Damit liegt am Ausgang nur die Sättigungsspannung von $V4$; das entspricht dem Pegel L. Ohne Berücksichtigung der Toleranzen ergeben sich die eingetragenen Potentiale und Ströme. An der Basis von $V4$ und damit am Emmitter von $V2$ liegt eine Span-

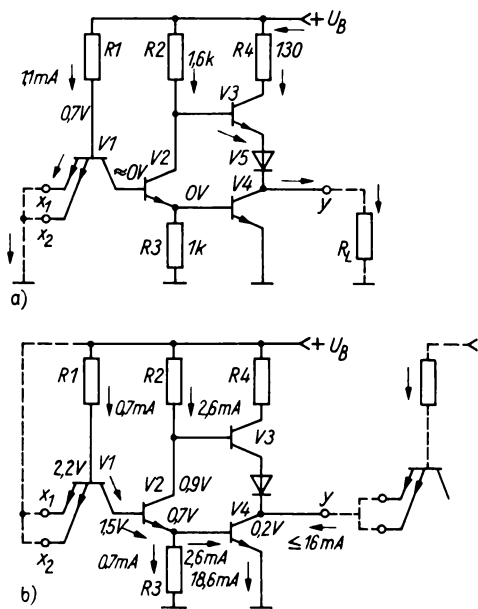


Bild 5.8. Schaltung eines TTL-Schaltkreises

a) bei Pegel L am Eingang; b) bei Pegel H am Eingang

nung von $0,7\text{ V}$. Die Spannung am Kollektor von $V2$ ist um $U_{CE\text{ sat}}$ größer als die Emitterspannung. Da zur Aufsteuerung von $V3$ mindestens ein Potential von $1,4\text{ V}$ an der Basis benötigt würde, jedoch nur $0,9\text{ V}$ vorhanden sind, ist $V3$ sicher gesperrt. Der Strom durch $R1$ ergibt sich durch die ebenfalls eingetragenen Potentiale. Bild 5.9 zeigt die Belastung eines TTL-Gatters durch weitere TTL-Gatter. Liegt am Ausgang Pegel L, so fließt in den Ausgang die Summe der Eingangsströme der nachgeschalteten Gatter hinein (Bild 5.9a). Der in den Kollektor von $V4$ hineinfließende Strom darf nicht zu groß werden, da bei dem eingespeisten konstanten Basisstrom sonst der Sättigungsbereich der Kennlinie verlassen würde. Die Folge wäre ein Anstieg der Ausgangsspannung und damit ein Verlassen des Pegelbereiches L. Die höchstzulässige Belastung wird durch den *Ausgangslastfaktor* N_o angegeben.

Beispiel

Für den integrierten Schaltkreis D100 ist $N_o = 10$ – der maximal aus dem Eingang eines Gatters der D10-Serie herausfließende Strom beträgt $1,6\text{ mA}$; daraus folgt, daß der Kollektorstrom von $V4$ maximal 16 mA betragen darf.)

Liegt am Ausgang Pegel H, so fließt aus dem Ausgang die Summe der Eingangsrestströme der folgenden Gatter heraus. Gemäß Bild 5.8b ist die Basis-Emitter-

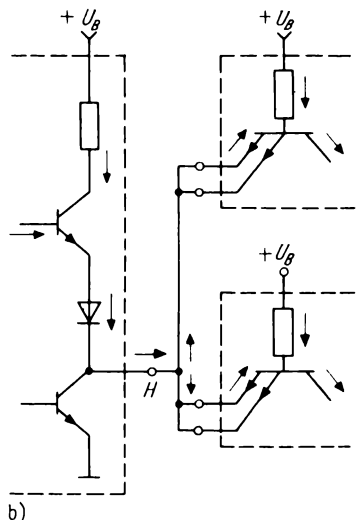
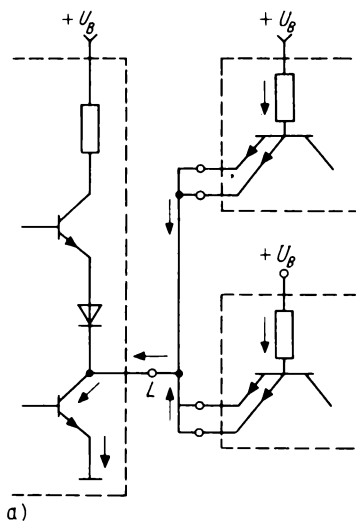


Bild 5.9. Belastung eines TTL-Gatters durch weitere TTL-Gatter

a) für Ausgangspegel L; b) für Ausgangspegel H

Diode der Eingangstransistoren dieser Gatter in Sperrichtung vorgespannt. Der maximal in den Gattereingang hineinfließende Sperrstrom wird bei der Serie D10 mit $40\ \mu\text{A}$ angegeben. Daraus folgt, daß aus einem Gatterausgang mit $N_o = 10$ maximal $0,4\ \text{mA}$ herausfließen dürfen.

Die notwendige Begrenzung dieses Stromes ist durch den von ihm verursachten Spannungsabfall an R_4 zu begründen. Mit steigendem Spannungsabfall sinkt der Ausgangspegel, um schließlich den zulässigen Bereich zu verlassen. Bild 5.10 zeigt die Abhängigkeit der Ausgangs-

spannung vom Laststrom für die vorstehend beschriebenen Fälle.

Offene Eingänge eines Gatters nach Bild 5.8 wirken wie auf H geschaltete Eingänge, da kein Strom fließen kann.

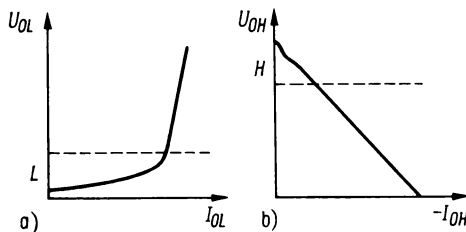


Bild 5.10. Abhängigkeit der Ausgangsspannung vom Ausgangsstrom

a) bei Ausgangspegel L; b) bei Ausgangspegel H

Der aus dem Eingang herausfließende Strom (s. Bild 5.8 a) wird von der wirksamen Spannung bestimmt und steigt mit negativer werdender Eingangsspannung. Um eine Überlastung des Eingangstransistors zu vermeiden, wird eine *minimal zulässige Eingangsspannung* als Grenzwert angegeben. (Zum Beispiel darf die Eingangsspannung im statischen Betrieb einen Wert von $-0,8\ \text{V}$ und bei definiertem dynamischem Betrieb von $-1,5\ \text{V}$ nicht unterschreiten.)

Nach Bild 5.8 b ist die Basis-Emitter-Diode des Eingangstransistors in Sperrichtung vorgespannt. Die zulässige Sperrspannung dieser Diode ist gering. Es wird deshalb eine *höchste Eingangsspannung* angegeben. (Ohne zusätzliche Strombegrenzungsschaltung darf die Eingangsspannung z. B. einen Wert von $5,5\ \text{V}$ nicht überschreiten.)

Mehrere Gatterausgänge nach Bild 5.8 dürfen nicht verbunden werden. L-Pegel am Ausgang des einen und H-Pegel am Ausgang des anderen würden einen Kurzschlußstrom zur Folge haben, da am ersten Gatter $V/4$ und am zweiten $V/3$ voll durchsteuert wären. Dieser Ausgang wird auch als „totem pole output“ bezeichnet, da er in beiden Zuständen niederohmig ist. Beim Ausgangspegel H spricht man deshalb auch von Zustand „pull up“ (der Ausgang hat eine niederohmige Verbindung zur Betriebsspannung) und bei Ausgangspegel L vom Zustand „pull down“ (der Ausgang hat eine niederohmige Verbindung zur Masse).

Um eine ausgangsseitige Parallelschaltung zu ermöglichen, werden Schaltkreise mit offenem Kollektor (engl., open collector output) angebo-

ten. Hierzu gehört beispielsweise der Schaltkreis D 103. Ein derartiger Schaltkreis ist im Bild 5.11 dargestellt. Man kann mehrere solcher Schaltkreisausgänge über einen externen Widerstand mit der Betriebsspannung verbinden. Es lassen sich so Montagelogikelemente zusammenstellen. Im Bild 5.12 sind zwei Ausgänge über einen gemeinsamen Widerstand mit der Betriebsspannung verbunden. Es liegt immer dann an y der Pegel L, wenn mindestens einer der beiden Ausgangstransistoren durchsteuert ist. Verbindet man jeweils x_1 und x_2 zu x_5 , sowie x_3 und x_4 zu x_6 , so erfüllen diese verbundenen Eingänge die Funktionsmatrix eines logischen Verknüpfungsglieds NOR. Durch Nachschaltung eines Inverters kann ein ODER realisiert werden.

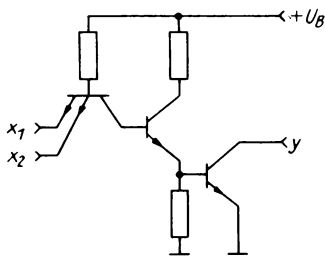


Bild 5.11. Schaltkreis mit offenem Kollektor

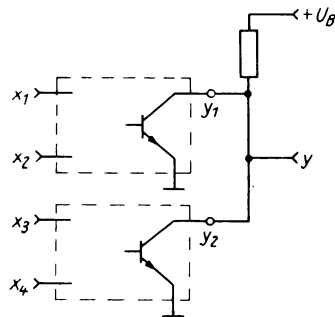


Bild 5.12. Verknüpfung zweier Schaltkreisausgänge

Es ist auch möglich, Gatterausgänge über eine gemeinsame Leitung (Sammelschienenübertragung, Bus) mit weiteren Schaltungsteilen zu verbinden. Selbstverständlich darf dann zu einem bestimmten Zeitpunkt nur ein Gatterausgang in die gemeinsame Leitung einspeisen; alle anderen Gatterausgänge müssen unabhängig von ihren Eingangszuständen am Ausgang hochohmig sein (scheinbar von der gemeinsamen Ausgangsleitung getrennt). Dieser dritte mögliche Ausgangszustand ist bei Gattern mit

tristablem Ausgang möglich (engl., three state output). Die Schaltung eines derartigen TTL-Gatters ist im Bild 5.13 dargestellt. Wird hier an den zusätzlichen Eingang S der Pegel H angelegt, so sind beide Ausgangstransistoren $V3$ und $V4$ gesperrt, und es fließen die eingezeichneten Ströme. Liegt der Pegel L an S , so hat der zusätzliche Schaltungsteil keinen Einfluß auf die Wirkungsweise des TTL-NAND.

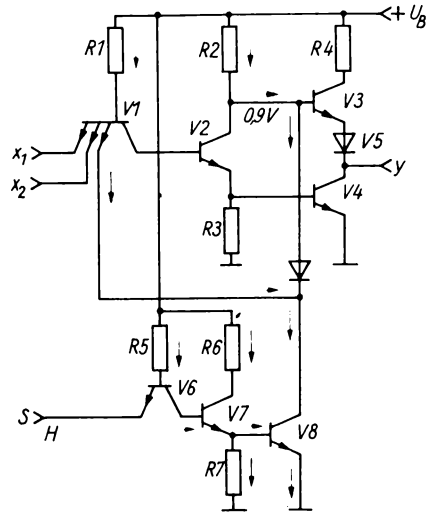


Bild 5.13. TTL-Gatter mit tristablem Ausgang

Die Schaltzeit des TTL-Gatters wird durch die Übersteuerung der Transistoren unerwünscht erhöht (beim Umschalten in den Sperrzustand müssen die Übersteuerungsladungen zusätzlich aus der Basiszone abgeführt werden). Andererseits ist die Übersteuerung für ein schnelles und vollständiges Durchsteuern des Transistors trotz fertigungsbedingter Toleranzen und unterschiedlicher Betriebstemperaturen erforderlich. Die Schaltung einer Schottky-Diode zwischen Kollektor und Basis (Bild 5.14) gestattet ein schnelles und sicheres Durchsteuern, ohne dabei den Transistor in die Sättigung zu treiben. Durch die niedrige Flußspannung der Schottky-Diode von etwa 0,4 V wird der zur Übersättigung der Basis führende Strom über die Diode abgeleitet und U_{CF} kann nur bis etwa 0,3 V sin-

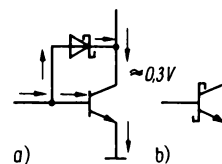


Bild 5.14
Transistor mit Schottky-Diode
a) Schaltung;
b) vereinfachte Darstellung

ken. Der Stromfluß ist für voll geöffneten Transistor auf Bild 5.14a eingetragen. Infolge der geringen Sperrerrholzeit der Diode ist ein schnelles Schalten in den Sperrzustand möglich. Bei der TTL-LS wird dieser Effekt genutzt, um bei etwa gleichbleibender Schaltzeit gegenüber TTL-Standard mit wesentlich geringerer Verlustleistung auszukommen.

Auf Bild 5.15 ist die Schaltung eines TTL-LS-NAND dargestellt. Im Eingang ist auf den Mehremittertransistor verzichtet. Gleichzeitig ist an jedem Eingang eine gegen Masse geschaltete Clamping-Diode ($V1$ und $V2$) zum Schutz vor negativen Spannungsspitzen integriert. (Diese Dioden sind nicht für permanente Belastung in Flußrichtung zugelassen.) Die gegenüber dem Mehremittertransistor höhere Spannungsfestigkeit der Eingangsdiolen ermöglicht als Grenzwert der Eingangsspannung ebenfalls 7 V wie als Grenzwert der Betriebsspannung (bei TTL-Standard darf U_I 5,5 V nicht überschreiten, wenn keine externen Strombegrenzungsmaßnahmen verwendet werden). Die weiteren Schaltungsänderungen gegenüber TTL-Standard dienen der zusätzlichen Schaltzeitverkürzung.

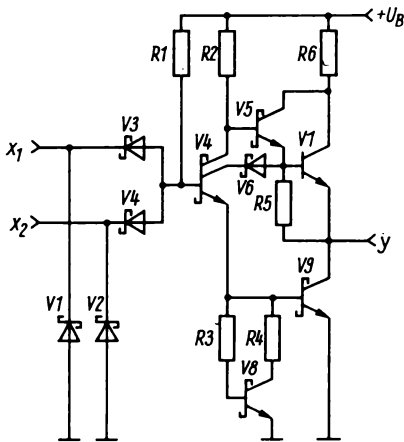


Bild 5.15. Schaltung eines TTL-LS NAND

Der Ausgangspegel unterscheidet sich geringfügig von dem des TTL-Standards. Es gilt: $U_{OL} < 0,5 \text{ V}$ und $U_{OH} > 2,7 \text{ V}$. Die Belastbarkeit des Normalausganges beträgt für Pegel L 8 mA und für H $-400 \mu\text{A}$, die Eingangsströme betragen maximal bei Pegel L $-0,4 \text{ mA}$ und bei Pegel H $20 \mu\text{A}$.

Die auf etwa $1/3$ reduzierte Verlustleistung bei sonst weitgehend gleichen Parametern führt zu

einer schrittweisen Ablösung des TTL-Standards durch die TTL-LS. In den meisten Schaltungen ist ein direkter Austausch der Schaltkreise ohne weitere Änderungen möglich. Im Sortiment der TTL-LS befinden sich auch AND- und NOR-Elemente.

In den vorangegangenen Abschnitten wurden der Eingangsspannung jeweils konstante Werte zugeordnet. Vielfach interessiert auch die Abhängigkeit der Ausgangsspannung von der Eingangsspannung. Die dazu im Bild 5.16 mit eingetragenen Pegelgrenzen dargestellte Kennlinie heißt *Transferkennlinie*. Die erkennbare zusätzliche Sicherheit zum Signalpegel ist erforderlich, um unter allen möglichen Toleranzen der Arbeitstemperatur, der Betriebsspannung und der Schaltkreise die definierten Pegel sicher einzuhalten.

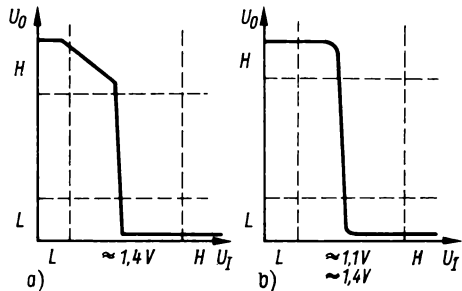


Bild 5.16. Transferkennlinie eines TTL-Schaltkreises

a) TTL-Standard; b) TTL-LS

Der Unterschied zwischen dem Pegel am Eingang und dem am Ausgang (s. Bild 5.1) ergibt die garantierte *statische Störsicherheit*. Im steil verlaufenden Teil der Transferkennlinie ergibt eine kleine Eingangsspannungsänderung eine große Ausgangsspannungsänderung. Das entspricht dem Verhalten eines Verstärkers mit hoher Verstärkung. Dieser steil verlaufende Kennlinienteil wird deshalb auch als „aktiver Bereich der Kennlinie“ bezeichnet und die mittlere Spannung dieses Teils als Schwellspannung. Im aktiven Bereich ist durch eine eventuell auftretende Mitkopplung die Gefahr der unerwünschten Selbsterregung groß. Um dies zu vermeiden, muß man dafür sorgen, daß die Eingangsspannung möglichst eine kurze Anstiegs- oder Abfallzeit hat. Der aktive Bereich wird dann schnell durchlaufen, und eine Selbsterregung kann nicht eintreten. Für die Serie D10 gilt beispielsweise, daß bei kombinatorisch wirkenden Schaltungen die Anstiegs- bzw. Abfallzeit kleiner als $1 \mu\text{s}$ sein muß. Ebenso wird gefordert,

daß die Lastkapazität an einem Ausgang einen Höchstwert nicht überschreiten darf (bei D10 1 nF), um unnötige Verzögerungen der Anstiegs- oder Abfallzeit der Ausgangsspannung zu verhindern, die ja meist als Eingangsspannung für weitere Gatter genutzt wird.

Um im dynamischen Betrieb Fehlschaltungen zu vermeiden, dürfen Signalleitungen nicht zu lang sein. (In der Praxis rechnet man mit Leitungslängen kleiner als 300 mm unter normalen Bedingungen.) Sind konstruktiv längere Leitungen erforderlich, müssen zusätzliche Schaltungsmaßnahmen getroffen werden.

Von Bedeutung für den praktischen Schaltungsaufbau ist auch die Abhängigkeit des Betriebsstroms von der Eingangsspannung (Bild 5.17).

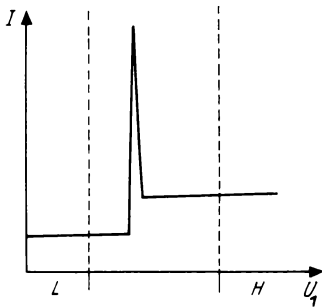


Bild 5.17. Abhängigkeit des Betriebsstroms von der Eingangsspannung

Die im Übergangsbereich von Eingangspegel L nach H liegende Stromspitze darf zu keinen Betriebsspannungsschwankungen führen, um eine Beeinflussung anderer Schaltkreise zu vermeiden. Es ist deshalb notwendig, die Betriebsspannung für etwa fünf bis zehn Schaltkreise durch einen Kondensator 10 ... 100 nF möglichst dicht an den Schaltkreisen und die Betriebsspannung für eine größere Gruppe durch einen Kondensator von etwa 10 μ F zur Ableitung von Störungen abzublocken (zusätzliche Herstellerangaben beachten!).

Der angegebene zulässige Betriebsspannungsbereich liegt zwischen 4,75 und 5,25 V ($5\text{ V} \pm 5\%$). Diese geringe zulässige Betriebsspannungstoleranz kann nur mit einer stabilisierten Stromversorgung erreicht werden. Der absolute Grenzwert für die Betriebsspannung beträgt 7 V. Bei Überschreitung dieses Grenzwerts ist eine Zerstörung der Schaltkreise möglich (s. Abschn. 1.4.3.).

Beim Eingangspegel L ist die Stromaufnahme niedriger als beim Eingangspegel H. Die Ein-

gänge unbenutzter Gatter eines Schaltkreises schaltet man deshalb zweckmäßig auf Masse. Unbenutzte Eingänge eines Gatters sollten ebenfalls nicht unbeschaltet bleiben. Nach Möglichkeit verbindet man sie mit bereits benutzten Eingängen oder schaltet sie auf Pegel H.

Beim Aufbau von gedruckten Schaltungen sollte stets für eine möglichst geringe Impedanz der Masseleitung gesorgt werden.

Zu beachten ist beim Aufbau der Schaltung auch die *Störkapazität* (bei TTL-LS etwa 37 pF, bei TTL-Standard 55 pF). Unter Störkapazität versteht man die sich zwischen der betrachteten Eingangsleitung und einer anderen Signalleitung bildende Kapazität, über die ein Pegelsprung in dieser Leitung einen unerwünschten Schaltvorgang im an die andere Leitung angeschlossenen Gatter auslösen kann. Die vom Hersteller gegebenen konkreten Meßbedingungen sind dabei zu beachten.

MOS-Gatter

Ein MOS-Grundgatter wird ausschließlich durch unipolare Transistoren realisiert. Es enthält im Unterschied zum TTL-Grundgatter folglich keine integrierten Widerstände. Der Flächenbedarf je Gatter wird damit wesentlich geringer. Da auch die erforderliche frequenzabhängige Verlustleistung weit unter der von TTL-Schaltkreisen liegt, sind höhere Integrationsgrade erzielbar. Nachteilig gegenüber TTL-Schaltkreisen sind die höheren Schaltzeiten.

Aus der Vielzahl möglicher MOS-Gatter sei zunächst das *p-Kanal-Hochvoltgatter* näher betrachtet. Die Gatter bestehen aus p-Kanal-Transistoren vom Anreicherungstyp (engl. p-channel enhancement type). Die Schaltung eines Inverters ist im Bild 5.18 dargestellt.

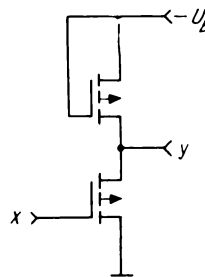


Bild 5.18
MOS-Inverter

Der Substratanschluß ist hierbei immer mit Masse verbunden. Die Aufgabe des Lastwiderstands übernimmt ein Transistor. Die Wir-

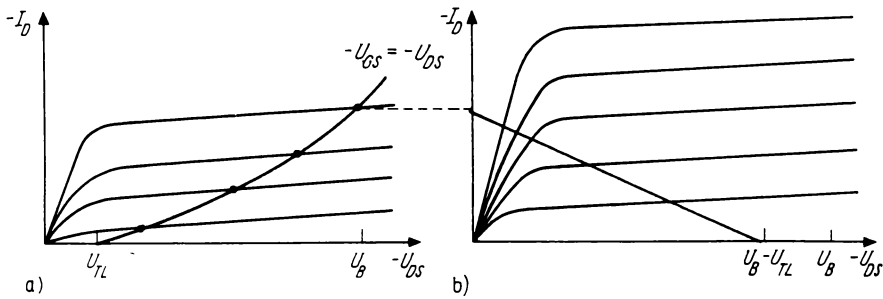


Bild 5.19. Konstruktion im Kennlinienfeld eines Inverters

a) Ausgangskennlinienfeld des Lasttransistors; b) Ausgangskennlinienfeld des Treibertransistors

kungsweise der Schaltung wird anhand der im Bild 5.19 dargestellten Ausgangskennlinien des Treiber- und des Lasttransistors erläutert. Dabei hat der Lasttransistor eine geringere Steilheit als der Treibertransistor. Entsprechend der Schaltung des Lasttransistors ist die sich ergebende Kennlinie für $-U_{GS} = -U_{DS}$ eingetragen. Aus dieser Kennlinie kann man den für die verschiedenen Drainströme am Lasttransistor auftretenden Spannungsabfall entnehmen. Um diesen Spannungsabfall bei dem jeweils beide Transistoren durchfließenden Drainstrom ist die Spannung am Drain des Treibertransistors positiver als die Betriebsspannung. Die punktweise Konstruktion liefert die „Widerstandsgerade“ im Kennlinienfeld des Treibertransistors. Sie unterscheidet sich nur wenig von der Widerstandsgeraden eines ohmschen Widerstands. Ist der Drainstrom 0, so ergibt sich der höchstmögliche Betrag der Ausgangsspannung des Inverters als um den Betrag der Schwellspannung des Lasttransistors erniedrigter Betrag der Betriebsspannung.

Die Transferkennlinie des Inverters wird im Bild 5.20 gezeigt. Bis zum Erreichen der Schwellspannung des Treibertransistors fließt

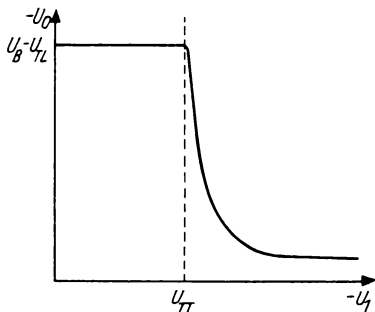


Bild 5.20. Transferkennlinie des MOS-Inverters

kein Strom durch den Inverter. Die Ausgangsspannung hat ihren höchsten Betrag. Nach Überschreiten der Schwellspannung wird der Betrag der Ausgangsspannung zunächst rasch kleiner, sinkt dann nach Erreichen des aktiven Bereichs des Treibertransistors etwas langsamer bis auf eine der jeweiligen Eingangsspannung entsprechende Spannung des Treibertransistors ab.

Die Transferkennlinie wurde für den unbelasteten Inverter betrachtet. Wird hingegen der Ausgang des Inverters bei Pegel L mit einem Widerstand belastet (Bild 5.21a), so fließt durch den Lasttransistor der eingezeichnete Strom. Aus Bild 5.19 kann man entnehmen, daß dabei die Ausgangsspannung infolge des Spannungsabfalls am Lasttransistor positiver wird. Liegt am Ausgang des Inverters Pegel H und wird er mit einem Widerstand gegen $-U_B$ belastet (Bild 5.21b), so fließt ein gegenüber dem unbelasteten Fall größerer Drainstrom durch den Treibertransistor, an dem sich der Spannungsabfall vergrößert. Die Ausgangsspannung wird negativer.

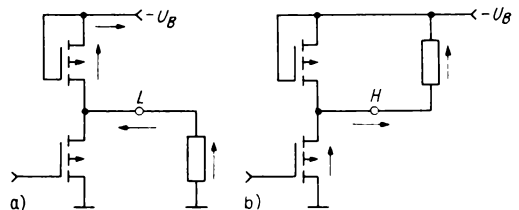


Bild 5.21. Belasteter Inverter

a) bei Ausgangspegel L; b) bei Ausgangspegel H

Die dargestellten Lastfälle können bei kapazitiver Last (Bild 5.22) auftreten. Die notwendigen Auf- und Entladeströme haben die im Bild 5.21 gezeigten Richtungen. Da die wirksamen Wi-

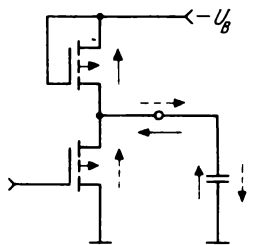
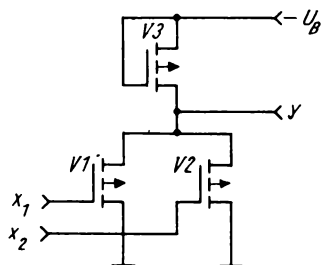


Bild 5.22. Inverter mit kapazitiver Last

derstände bei der Umladung der Kondensatoren groß sind, ergeben sich hohe Schaltzeiten.

Die durch die Lastströme auftretende Ausgangsspannungsänderung ist bei Belastung der Gatter zu beachten. Vom Hersteller werden die Pegelverschiebungen bei bestimmten Lastströmen (z. B. ± 1 mA) angegeben. Ähnlich der Bedingungen bei TTL-Gattern darf die Eingangsspannung nur in einem bestimmten Bereich liegen. Die zum Schutz gegen statische Aufladungen eingefügten Gateschutzdioden würden bei Überschreitung des Eingangsspannungsbereichs zerstört werden. Besonders der Pegel H darf höchstens einen gering positiven Wert annehmen.

Die beschriebene Inverterschaltung läßt sich leicht zu weiteren logischen Verknüpfungsgliedern erweitern. Bild 5.23 zeigt ein **NOR-Gatter**.



x_1	x_2	Y
H	H	L
L	H	H
H	L	H
L	L	H

Pegeltafel

x_1	x_2	Y
0	0	1
1	0	0
0	1	0
1	1	0

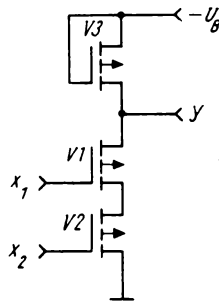
Wahrheitstabelle

Bild 5.23

MOS NOR-Gatter

Es wird immer dann am Ausgang Pegel H liegen, wenn am Eingang mindestens einer der beiden Transistoren V1 oder V2 Pegel L liegt. Da der Schaltkreis für negative Logik vorgesehen ist, folgt aus der im Bild 5.23 dargestellten Pegeltafel die zugehörige Wahrheitstabelle. In ähnlich einfacher Form läßt sich auch ein **NAND-Gatter** aufbauen (Bild 5.24). Hier kann

nur dann am Ausgang Pegel H liegen, wenn beide Transistoren V1 und V2 am Eingang auf Pegel L geschaltet sind. Daraus folgen wiederum Pegeltafel und Wahrheitstabelle. (Würde man die Pegel der positiven Logik zuordnen, so würde jeweils aus dem NAND ein NOR und aus dem NOR ein NAND.)



x_1	x_2	Y
H	H	L
L	H	L
H	L	L
L	L	H

Pegeltafel

x_1	x_2	Y
0	0	1
1	0	1
0	1	1
1	1	0

Wahrheitstabelle

Bild 5.24. MOS NAND-Gatter

MOS-Schaltkreise sind spannungsgesteuert; es fließt nur ein sehr geringer Reststrom als Eingangsstrom. An unbeschalteten Eingängen sind statische Aufladungen möglich, die Fehlschaltungen verursachen könnten. Nicht benötigte Eingänge muß man deshalb entsprechend der logischen Funktion der Schaltung entweder mit bereits benutzten Eingängen verbinden, an Pegel H (Masse) oder an möglichst nur minimalen Pegel L schalten (über Spannungsteiler an $-U_B$). Da die Eingänge hochohmig sind, genügen bereits geringe **Störkapazitäten** (etwa 25 pF), um auf den Eingang einen Störimpuls zu koppeln. Die Verlegung der Eingangsleitungen ist deshalb besonders sorgfältig auszuführen. Für die Abblockung der Betriebsspannung gilt das bei TTL-Schaltkreisen Gesagte. Die integrierten Gateschutzdioden sind nicht zu schaltungstechnischen Zwecken (z. B. als Begrenzer) zu benutzen.

Einen noch niedrigeren Ruheleistungsverbrauch als die vorstehend beschriebenen Gatter haben **CMOS-Schaltungen**, bei denen sowohl n-Kanal- als auch p-Kanal-Transistoren verwendet werden. Die Herstellung dieser Gatter ist jedoch schwieriger als die der vorstehend beschriebenen Gatter. Ein CMOS-Inverter ist im Bild 5.25 dargestellt. Er erhält eine positive Betriebsspannung; sein Signalpegel sei der positiven Logik zugeordnet. Liegt am Eingang Pegel L, so sind der Treibertransistor gesperrt und der Lasttransistor geöffnet. Am Ausgang liegt Pegel H. Schaltet man hingegen den Eingang auf Pe-

gel H, so sind der Treibertransistor voll durchsteuert und der Lasttransistor gesperrt. Am Ausgang liegt Pegel L. Für jeden der beiden möglichen Schaltzustände ist ein Transistor gesperrt. Daraus ergibt sich die geringe Ruheverlustleistung. Wegen der erforderlichen Umladung der strukturbedingten Kapazitäten steigt die Verlustleistung mit höher werdender Frequenz.

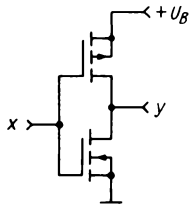


Bild 5.25
CMOS-Inverter

Der Ausgangspegel entspricht beim Pegel H nahezu der Betriebsspannung (pull up), da die Flußspannung des dabei leitenden Lasttransistors nur wenige Millivolt beträgt, und beim Pegel L wiederum wegen der niedrigen Flußspannung des jetzt durchsteuerten Treibertransistors nur wenige Millivolt (pull down). Der Ausgangspegelhub erreicht damit fast den Wert der Betriebsspannung. Der in beiden Fällen niederohmige Ausgang entspricht damit dem „totem pole output“ bei TTL-Schaltkreisen.

Eine CMOS-NOR-Schaltung wird im Bild 5.26 gezeigt. Liegt hier an einem der beiden Eingänge H, so sind der zugehörige Treibertransistor geöffnet und der Lasttransistor gesperrt. Da die beiden Treibertransistoren parallel und die Lasttransistoren in Reihe geschaltet sind, kann sich am Ausgang nur der Pegel L einstellen. Nur wenn beide Eingänge auf Pegel L geschaltet sind, sind die Treibertransistoren gesperrt und die Lasttransistoren geöffnet. Am Ausgang erscheint Pegel H.

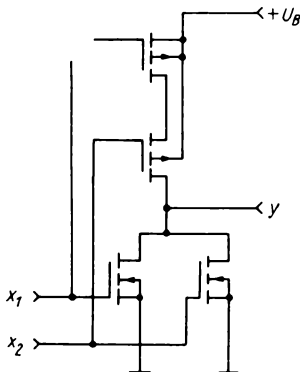


Bild 5.26
CMOS-NOR

In ähnlicher Form ist auch das CMOS-NAND-Gatter aufgebaut (Bild 5.27). Hier sind jeweils die Treibertransistoren in Reihe und die Lasttransistoren parallelgeschaltet. Am Ausgang ist nur dann Pegel L möglich, wenn beide Treibertransistoren durch Pegel H geöffnet sind. Bereits wenn an einem Eingang Pegel L liegt, ist die Reihenschaltung der Treibertransistoren gesperrt, dafür mindestens einer der Lasttransistoren geöffnet. Am Ausgang kann nur Pegel H liegen.

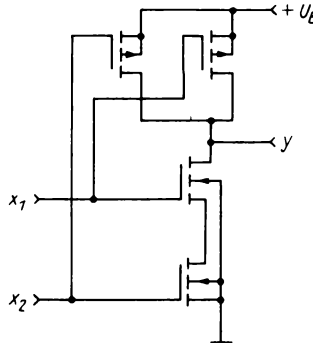


Bild 5.27
CMOS-NAND

In den gezeigten Schaltungen besteht eine nachteilige Rückwirkung des Ausgangs auf den Eingang. Die meisten heute gefertigten Schaltkreise der Standardreihe erhalten deshalb einen Ausgangspuffer (engl., buffered), bestehend aus meist zwei hintereinander geschalteten Invertieren (Bild 5.28). Trotz der zwei nachgeschalteten Stufen ergibt sich keine Erhöhung der Schaltzeit, da jetzt die Transistoren der logischen Verknüpfungen für wesentlich kleinere Ströme ausgelegt werden können und damit durch schnelleres Schalten die zusätzliche Schaltzeit des Ausgangspuffers kompensieren. Alle Eingänge erhalten integrierte Gateschutzschaltungen, die unterschiedlich aufgebaut sein

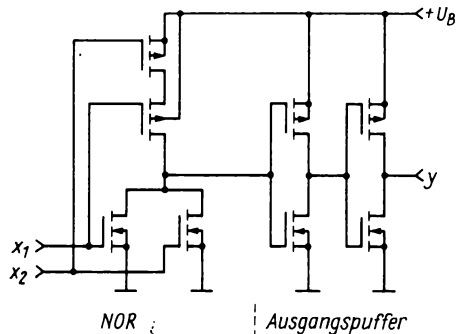


Bild 5.28. CMOS-NOR mit Ausgangspuffer

können. Eine Variante zeigt Bild 5.29. Die beiden Dioden schützen sicher sowohl vor positiven als auch vor negativen statischen Aufladungen. Der in Reihe liegende Widerstand dient zur Strombegrenzung durch die Dioden. Die Dioden sind für die Ableitung von Aufladungen bemessen, sie dürfen nicht für schaltungstechnische Zwecke genutzt werden.

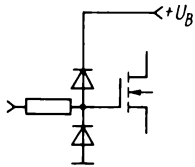


Bild 5.29
Gateschutzschaltung

Es ist stets darauf zu achten, daß Eingangsspannungen nur dann anliegen können, wenn auch die Betriebsspannung anliegt (zu beachten bei Kopplung von aus unterschiedlichen Stromversorgungsteilen gespeisten Schaltungsteilen), da beim Wegfall der Betriebsspannung und Pegel H am Eingang es wegen der dann nicht mehr vorgespannten oberen Diode zu einem diese Diode zerstörenden Stromfluß kommen könnte. Schaltkreise mit open collector output werden bei CMOS nicht gefertigt. Hingegen sind Schaltkreise mit three state output im Angebot. Sie werden – wie auch bei TTL – durch einen zusätzlichen Eingang hochohmig geschaltet oder aktiviert.

Die *Transferkennlinie* (Bild 5.30) hat fast den idealen rechteckigen Verlauf bei gepuffertem Ausgang. Die Schwellspannung liegt etwa bei $0,6 U_B$, sie ist von der Betriebsspannung abhängig. Der Pegelhub entspricht praktisch der Betriebsspannung.

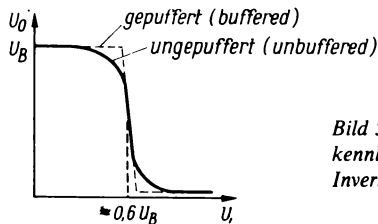


Bild 5.30. Transferkennlinie des CMOS-Inverters

Im Gegensatz zu TTL benötigen CMOS-Schaltkreise keine stabilisierte Betriebsspannung. Sie sind von 3 V bis 15 V funktionsfähig, wobei die Schaltzeit mit sinkender Betriebsspannung steigt. Für die Abblockung der Betriebsspannung gelten die bereits bei TTL gegebenen Ausführungen.

Schaltungstechnische Maßnahmen erfordert der *latch-up-Effekt*. Er kann entstehen, wenn an einem Anschluß des Schaltkreises die Spannung negativer als $-0,5 V$ oder positiver als $U_B + 0,5 V$ wird bzw. der Strom 10 mA übersteigt. Unter diesen Bedingungen ist das Zünden einer technologiebedingten, durch parasitäre pn-Übergänge gebildeten Thyristorstruktur möglich. Fließt dabei ein Strom größer als der Haltestrom von 10 mA, kommt es zur Zerstörung des Schaltkreises. Man muß deshalb entweder das Auftreten der unzulässigen Spannungen sicher verhindern oder durch Strombegrenzungswiderstände den möglichen Stromfluß auf kleiner als 10 mA begrenzen.

Offene Eingänge sind bei CMOS grundsätzlich nicht zulässig. Infolge der Hochohmigkeit würde sich ein beliebiges Potential aufbauen können mit der Folge, daß Fehlschaltungen auftreten würden. Nicht benutzte Eingänge sind deshalb über einen Widerstand von 100 k Ω bis 1 M Ω je nach der logischen Funktion auf Masse oder Betriebsspannung zu schalten. Diese Maßnahme ist auch für an Steckverbinder angeschlossene Eingänge empfehlenswert. Ebenfalls sollten die von three state output gespeisten Busleitungen so beschaltet werden.

Die Abhängigkeit der Ausgangsspannung vom Ausgangsstrom zeigt Bild 5.31. Es herrschen ähnliche Verhältnisse wie bei TTL. Auf konkrete Werte wird im Abschnitt 5.3.3. eingegangen.

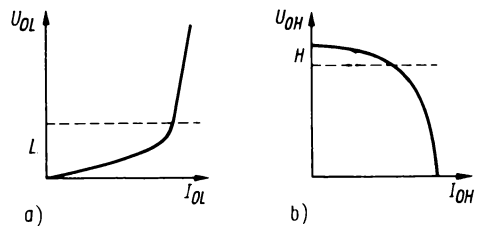


Bild 5.31. Abhängigkeit der Ausgangsspannung vom Ausgangsstrom

a) bei Ausgangspegel L; b) bei Ausgangspegel H

Gegenüber TTL ist bis zu mittleren Verarbeitungsfrequenzen die Verlustleistung wesentlich kleiner. Besser ist auch der statische Störabstand sowie die äußerst geringe Temperaturabhängigkeit der Schwellspannung und der große Umgebungstemperaturbereich (bei CMOS -40 bis $+85 ^\circ C$, bei TTL mit erweitertem Temperaturbereich -25 bis $+85 ^\circ C$). Nachteilig hinge-

gen wirkt sich die heute noch geringere Fertigungsausbeute, größere Schaltzeit und die gegenüber der langjährig eingesetzten TTL geringere Einsatzerfahrung aus. Außerdem ist der CMOS-Standard nicht pin-kompatibel zu TTL. Zur Überwindung obiger Nachteile wurde in jüngster Zeit die CMOS-HCT entwickelt. Es handelt sich dabei um eine schnelle CMOS, deren besonderer Vorteil in der Kompatibilität zu TTL liegt. So entspricht die pin-Belegung voll der von TTL-Schaltkreisen. Die Pegel von CMOS-HCT, die Ausgangsbelastbarkeit sowie Schaltzeit gestatten eine direkte Verknüpfung mit TTL (siehe Tafeln 5.2 und 5.3). Da die sonstigen Vorteile von CMOS gegenüber TTL erhalten sind, ist eine künftige Ablösung von TTL durch CMOS-HCT nicht nur bei Neuentwicklungen sondern auch bei Erweiterung oder Ersatzbestückung bestehender Geräte denkbar. Bei der CMOS-HCT liegt international derzeit die größte Steigerungsrate der Produktion vor. Im Gegensatz zu bipolaren Schaltkreisen haben unipolare Schaltkreise eine hohe Strahlungsresistenz. Sie sind deshalb besonders für Anlagen der Landesverteidigung und der Raumfahrt geeignet.

ECL-Gatter

Die in den vorangegangenen Abschnitten beschriebenen Gatter gehören zur Übersteuerungstechnik. Relativ große Toleranzen der Bauelementewerte bleiben dabei wirkungslos auf die Funktion des Schaltkreises. Die Übersteuerung verursacht jedoch höhere Schaltzeiten, da die Übersteuerungsladungen jeweils erst abfließen müssen. Die emittergekoppelte Logik (ECL) beruht auf dem *Stromsteuerprinzip* und ist keine Übersteuerungslogik. Schaltzeiten in der Größenordnung von nur 1 ns sind möglich. Die Herstellung ist kompliziert, da besonders die Widerstände mit niedriger zulässiger Toleranz integriert werden müssen. Da auch ständig Strom fließt, ist die Verlustleistung je Gatter hoch. Der Einsatz der ECL-Gatter ist deshalb nur dort vertretbar, wo extrem niedrige Schaltzeiten gefordert sind.

Im Bild 5.32 ist ein ECL-Inverter dargestellt. Zum besseren Verständnis sind mögliche Widerstandswerte und Spannungen eingetragen. Liegt an der Basis von $V1$ eine gegenüber der Referenzspannung U_{Ref} kleinere Spannung, so wird nur durch $V2$ Strom fließen. Da die Transistoren nicht im Übersteuerungsbereich arbeiten sollen, sind die Widerstandswerte so gewählt,

daß an $\bar{y}$ eine etwas über der Referenzspannung liegende Spannung und an y die Betriebsspannung liegt. Die Spannung am Emitter ist um die Diodenflußspannung niedriger als die Referenzspannung. Daraus folgt der mögliche Emittierstrom (im Beispiel etwa 0,83 mA). Da der Kollektorstrom von $V2$ etwa gleich dem Emittierstrom ist, kann der Spannungsabfall am Lastwiderstand ermittelt werden (hier etwa 0,83 V) und damit die Spannung an $\bar{y}$ (hier etwa 4,17 V). Legt man an die Basis von $V1$ eine Spannung gleich der Referenzspannung, so teilt sich der Strom auf beide Transistoren auf. Eine Gesamtstromerhöhung ist durch die gegenkoppelnde Wirkung des Emittierwiderstands nicht möglich. Bei gegenüber der Referenzspannung größerer Spannung an der Basis von $V1$ erfolgt die Stromübernahme auf $V1$. Damit liegen an $\bar{y}$ die Betriebsspannung und an y eine gegenüber der Referenzspannung geringfügig höhere Spannung.

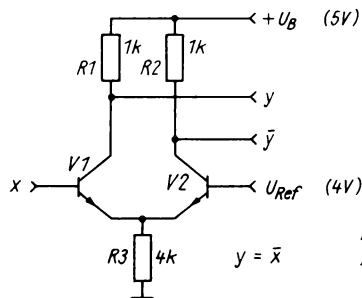


Bild 5.32.
ECL-Inverter

Mit den im Bild 5.32 angenommenen Spannungswerten ergibt sich eine Transferkennlinie nach Bild 5.33. Übersteigt die Eingangsspannung einen bestimmten Wert, so gelangt $V1$ in den Sättigungsbereich, die Ausgangsspannung steigt. Dieser Bereich muß außerhalb der normalen Betriebsbedingungen liegen. Aus den

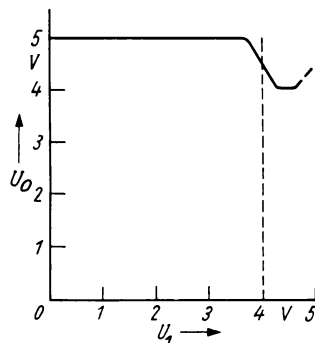


Bild 5.33. Transferkennlinie eines ECL-Inverters

vorangegangenen Darlegungen und den Werten im Bild 5.33 erkennt man, daß der Ausgangspegel nicht mit dem Eingangspegel kompatibel ist. Eine Pegelanpassung nach Bild 5.34 ist erforderlich. Die nachgeschalteten Emitterfolgerstufen senken den Ausgangspegel um eine Diodenflußspannung ab. (Damit durch die nachgeschalteten Transistoren Strom fließen kann, muß die Basis um etwa 0,7 V positiver als der Emitter sein.) Die vorgegebenen Zahlenwerte lassen auch erkennen, daß der erreichbare Pegelhub klein ist.

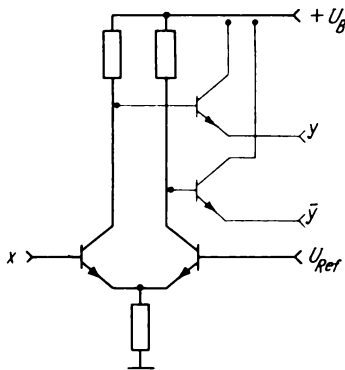


Bild 5.34. ECL-Inverter mit Pegelanpaßstufe

Soll ein NOR-Gatter realisiert werden, so schaltet man parallel zu V1 einen weiteren Transistor V3 (Bild 5.35). Durch die Parallelschaltung ist es gleichgültig, welcher der beiden Transistoren V1 oder V3 aufgesteuert wird, um an y Pegel L zu erhalten. Nur wenn an x₁ und x₂ Pegel L liegt, kann an y Pegel H entstehen. Die benötigte Referenzspannung kann man von außen zuführen oder durch eine interne Schaltung im Schaltkreis selbst erzeugen.

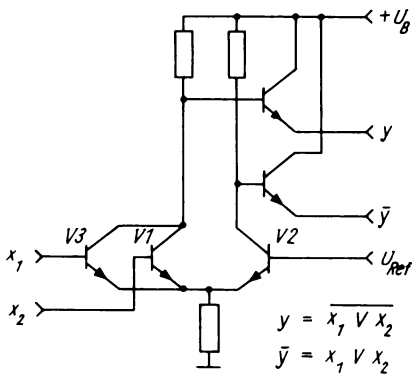


Bild 5.35. ECL-NOR/OR-Gatter

I²L-Gatter

Eine besonders hohe Packungsdichte bei geringer Anzahl von Herstellungsschritten, geringe Verlustleistung und hohe Schaltgeschwindigkeit kennzeichnen die noch junge I²L-Technik (integrierte Injektionslogik). Für einen hohen Integrationsgrad ist sie besonders gut geeignet. Durch Variation des Injektorstroms ist der Einsatz sowohl bei hoher Schaltgeschwindigkeit als auch bei besonders niedriger Verlustleistung möglich.

Die sonst übliche Betriebsspannung wird durch einen *Injektor* ersetzt, der einen definierten Strom in die Schaltung einspeist. Im Bild 5.36 ist ein I²L-Inverter dargestellt, dessen Eingang an den Ausgang eines vorgeschalteten und dessen Ausgang an den Eingang eines nachgeschalteten I²L-Gatters angeschlossen ist. Bild 5.36a zeigt den möglichen Stromfluß für gesättigten Transistor V3, über den der durch V1 injizierte Strom abfließt. An x liegt somit nur die Sättigungsspannung von V3 (Pegel L), damit ist V2 gesperrt. Der durch V4 injizierte Strom fließt über V5 ab, somit liegt an y eine Diodenflußspannung (entspricht Pegel H). Ist V3 hingegen gesperrt (Bild 5.36b), so fließt der durch V1 injizierte Strom durch V2 und sättigt diesen. An x liegt eine Diodenflußspannung. Da V2 gesättigt

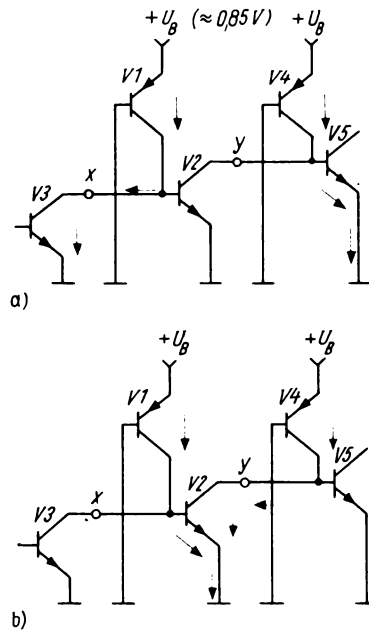


Bild 5.36. I²L-Inverter

a) Pegel L am Eingang; b) Pegel H am Eingang

ist, fließt der durch $V4$ injizierte Strom über $V2$ ab, und an y liegt Pegel L. Zur Vereinfachung zeichnet man anstelle der Injektoren eine Stromquelle. Vielfach werden die Transistoren auch mit mehreren unabhängigen Kollektoren hergestellt (Bild 5.37). Ist der Transistor gesättigt, so sind alle Kollektoren gesättigt; für gesperrten Transistor sind die Potentiale der Kollektoren voneinander unabhängig. Für verschiedene Verknüpfungsschaltungen und Speicherschaltungen ist dies vorteilhaft.

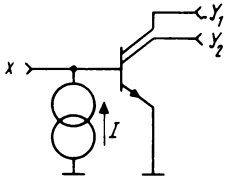


Bild 5.37. Vereinfachte Darstellung eines I^2L -Inverters mit mehreren Ausgängen

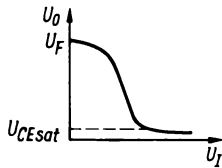


Bild 5.38. Transferkennlinie des I^2L -Inverters

Die Transferkennlinie ist im Bild 5.38 dargestellt. Die Ausgangsspannung liegt zwischen U_F (Flußspannung einer Diode) und U_{CEsat} (Kollektor-Emitter-Sättigungsspannung). Man erkennt, daß der logische Hub etwas kleiner als die Injektorspannung ist.

Die Darstellung eines NOR/OR-Gatters im Bild 5.39 zeigt nochmals die einfache Struktur von I^2L -Schaltungen. Ist mindestens einer der beiden Transistoren $V1$ oder $V2$ gesättigt (durch Pegel H am Eingang), so fließt der in den Ausgang eingespeiste Strom durch diesen Transistor ab. Am Ausgang y liegt Pegel L. Nur wenn beide Eingangstransistoren gesperrt sind (durch Pegel L am Eingang), gelangt der Ausgang y auf Pegel H. Die Schaltkreise sind der *positiven* Lo-

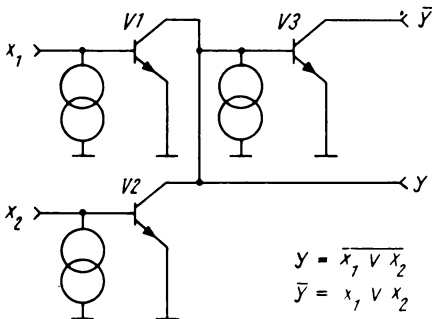


Bild 5.39. I^2L -NOR/OR-Gatter

gik zugeordnet. Der Ausgang y erfüllt für die Eingänge x_1 und x_2 die Funktionsmatrix eines NOR-Gatters. $V3$ negiert das an y liegende Signal; somit erfüllt $\bar{y}$ die Bedingung eines OR-Gatters für die Eingänge x_1 und x_2 . Entsprechend ihrer Eigenschaften wird die I^2L -Technik ausschließlich innerhalb hochintegrierter Schaltungen verwendet. So werden z. B. Teilerschaltkreise hoher Teilerfaktoren (D 351, D 355) oder Analog-Digital-Wandler (C 520) in dieser Technik realisiert, wobei Ein- und Ausgänge meist TTL-kompatibel gestaltet werden.

5.3.3. Interface-Schaltungen

In der Praxis kommt es häufig vor, daß in einem Gerät Schaltkreise aus mehreren Schaltkreisfamilien verwendet werden sollen. Aus den vorangegangenen Abschnitten folgt, daß jede Schaltkreisfamilie einen eigenen Signalpegel hat – die Pegel sind nicht kompatibel. Zur Anpassung der unterschiedlichen Pegel zweier Schaltkreisfamilien werden Interface-Schaltungen eingesetzt. Von Interface-Schaltungen spricht man auch, wenn an einem Schaltkreis systemfremde Lasten anzuschließen sind.

Die Notwendigkeit zum Einsatz von Schaltkreisen mehrerer Schaltkreisfamilien kann bestehen, wenn in einem mit TTL-Schaltkreisen aufgebauten Gerät LSI-Schaltkreise, wie Speicher hoher Kapazität, Teiler mit sehr hohen Teilerverhältnissen oder Mikroprozessoren benötigt werden, die meist in MOS- oder I^2L -Technik realisiert sind. Vielfach sind in diesen LSI-Schaltkreisen entsprechende Interface-Schaltungen bereits mit integriert. Am Ausgang findet dabei häufig eine Stufe mit offenem Kollektor (engl., open collector output) Verwendung. Als Beispiel seien die I^2L -Schaltkreise E 350D und E 355D genannt (Teilerschaltkreise mit sehr hohen Teilerfaktoren), deren Eingänge und Ausgänge durch integrierte Interface-Schaltungen TTL-kompatibel sind und so eine bequeme Anpassung an die Peripherie ermöglichen.

Bei der konkreten Prüfung der Zusammenschaltbarkeit muß man neben dem Pegel unbedingt die Ein- und Ausgangsströme beachten. In Tafel 5.3 sind deshalb für das Grundgatter der vier am häufigsten eingesetzten Schaltkreisfamilien die notwendigen Daten zusammengestellt. Man erkennt, daß z. B. der direkte Anschluß von CMOS an TTL wegen der nicht ausreichenden U_{OH} von TTL und der Anschluß von

Tafel 5.3. Ein- und Ausgangswerte wichtiger Schaltkreisfamilien

	U_{OH} bei $-I_{OH}$		U_{OL} bei I_{OL}		U_{IH}	I_{IH}	U_{IL}	$-I_{IL}$
TTL-Stand.	2,4 V	0,4 mA	0,4 V	16 mA	2 V	40 μ A	0,8 V	1,6 mA
TTL-LS	2,7 V	0,4 mA	0,5 V	8 mA	2 V	20 μ A	0,8 V	0,36 mA
CMOS-Stand. bei $U_B = 5$ V	4,6 V	0,4 mA	0,4 V	0,4 mA	3,5 V	1 μ A	1,5 V	1 μ A
	4,95 V	1 μ A	0,05 V	1 μ A				
CMOS-HCT bei $U_B = 5$ V	4,2 V	5 mA	0,4 V	5 mA	2 V	1 μ A	0,8 V	1 μ A
	4,9 V	1 μ A	0,1 V	1 μ A				

TTL-Standard an CMOS wegen des hohen Eingangsstromes bei Pegel L von TTL nicht möglich sind. Hingegen kann TTL-LS direkt an einen CMOS-Ausgang angeschlossen werden.

Durch einen zusätzlichen *pull up-Widerstand* (1,5...4,7 k Ω) kann die U_{OH} von TTL erhöht werden, um damit den direkten Anschluß von CMOS zu ermöglichen. Bild 5.40 zeigt diese Schaltung. U_{OH} wird, da nur ein äußerst kleiner Reststrom durch R fließt, praktisch den Wert 5 V annehmen. Eine direkte Zusammenschaltbarkeit besteht zwischen TTL und CMOS-HCT. An zwei weiteren Beispielen sollen Interface-Schaltungen in diskreter Technik erläutert werden.

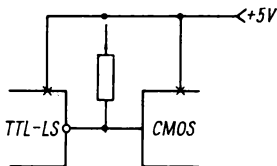


Bild 5.40. Pull-up-Widerstand zur direkten Kopplung von CMOS an TTL-LS (R : 1,5 ... 4,7 k Ω)

den. Bild 5.41 zeigt eine Möglichkeit zur Anpassung eines MOS-Ausgangs an einen TTL-Eingang. Da der MOS-Schaltkreis zur negativen Logik gehört, ist eine Negation der Pegel durch den Transistor erforderlich. Es muß z. B. Pegel L am Ausgang des MOS-Schaltkreises auf Pegel H am Eingang des TTL-Schaltkreises umgesetzt werden. Dabei wird die hohe negative Spannung am Ausgang des MOS-Schaltkreises über den Basisspannungsteiler so geteilt, daß der Transistor gesperrt bleibt. Am Kollektor des Transistors liegt dann eine Spannung von 5 V. Hat hingegen der Ausgang des MOS-Schaltkreises Pegel H (geringe negative Spannung), so

wird durch den Spannungsteiler der Transistor aufgesteuert, und am Kollektor liegt die Sättigungsspannung (Pegel L bei TTL). Bild 5.42 zeigt die Anpassung eines TTL-Ausgangs an einen MOS-Eingang. Auch hier ist eine Negation erforderlich. Der Basisspannungsteiler ist wiederum so bemessen, daß bei Eingangspegel H der Transistor gesperrt und bei Eingangspegel L der Transistor geöffnet ist.

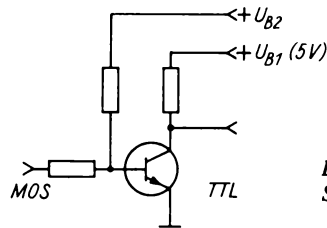


Bild 5.41. Interface-Schaltung MOS-TTL

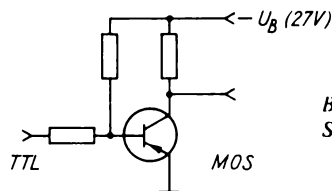


Bild 5.42. Interface-Schaltung TTL-MOS

Zusammenfassung

Logische Verknüpfungsglieder lassen sich in verschiedenen Techniken aufbauen. Zur Gruppe der Übersteuerungstechnik gehören die TTL-, die MOS- und die I²L-Gatter. Das Stromsteuerprinzip findet beim ECL-Gatter Anwendung. Die Schaltkreisfamilien unterscheiden sich insbesondere durch den technologischen Aufwand bei der Herstellung, die erzielbare Packungsdichte, die mögliche Schaltgeschwindigkeit, die Verlustleistung, die Pegel und die Betriebsspannung. Es sind jeweils die besonde-

ren Zusammenschaltungsbedingungen und die zulässigen Grenzwerte zu beachten. Durch die Transferkennlinie wird der statische Störabstand beschrieben. Sollen in einem Gerät Schaltkreise mehrerer Familien Verwendung finden, sind Interface-Schaltungen erforderlich. Diese beziehen sich häufig auf den TTL-Pegel.

5.4. Sequentielle Schaltungen

5.4.1. Allgemeines

Im Gegensatz zu kombinatorischen Schaltungen sind bei sequentiellen Schaltungen Rückführungen von Ausgängen auf Eingänge vorhanden. Es gibt zwei Signalflußrichtungen. Der Pegel am Ausgang ist vom vorherigen Zustand und vom Eingangspegel abhängig. Die wichtigsten Vertreter sind die Flip-Flop (auch Trigger oder bistabile Kippstufen genannt), aus deren Vielfalt der Varianten drei typische Vertreter ausgewählt worden sind. Signalgeneratoren, monostabile Glieder und Schwellwertelemente sind weitere Schaltungen dieser Gruppe.

5.4.2. RS-Flip-Flop

Der RS-Flip-Flop stellt die einfachste Form eines Flip-Flop dar. Er hat, wie alle Flip-Flop, zwei zueinander negierte Ausgänge, mit Q und $\bar{Q}$ bezeichnet (Bild 5.43). Gelangt der Eingang S (set, setzen) in den Einszustand, so erscheint am Ausgang Q ebenfalls der Einszustand. Der Ausgang Q behält dieses Signal auch dann, wenn S wieder den Nullzustand annimmt. Der Flip-Flop hat das Eingangssignal gespeichert.

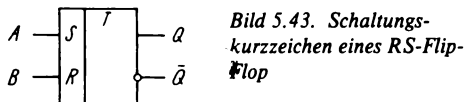


Bild 5.43. Schaltungskurzzeichen eines RS-Flip-Flop

Erst wenn an den Eingang R (reset, rücksetzen) ein Signal mit dem logischen Wert 1 gelegt wird, schaltet der Ausgang Q zurück auf den Nullzustand. Auch dieser bleibt erhalten (gespeichert), wenn R wieder in den Nullzustand gelangt.

Bild 5.44 zeigt die logische Struktur eines aus zwei NOR-Elementen aufgebauten RS-Flip-

Flop. Aus der Wahrheitstabelle des logischen Verknüpfungsgliedes NOR folgt, daß bereits bei Einszustand an einem Eingang der Ausgang Nullzustand annimmt.

Wird nun an S (entspricht Eingang A im Bild 5.44) ein Signal mit dem logischen Wert 1 gelegt, erscheint an $\bar{Q}$ der Nullzustand. Dieser logische Wert 0 wird an den Eingang des Gatters $D1$ geführt, an dessen Ausgang bei gleichzeitigem Nullzustand des Eingangs R (entspricht Eingang B im Bild 5.44) Einszustand entsteht. Da dieser Ausgang mit einem Eingang des Gatters $D2$ verbunden ist, bleibt auch nach Wegfall des Einszustand am Eingang S ein Eingang des Gatters $D2$ im Einszustand. An Q ist der Einszustand gespeichert. Der beschriebene Zustand ist im Bild 5.44 eingetragen. Infolge der Symmetrie der Schaltung laufen die Vorgänge bei Einszustand an R entsprechend ab.

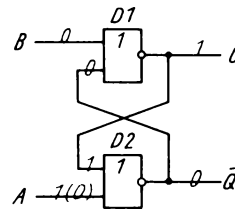
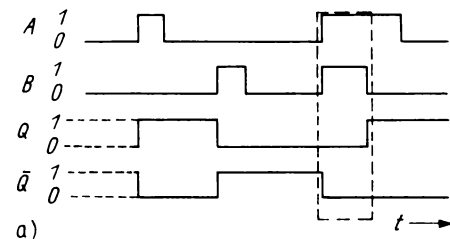


Bild 5.44. Logische Struktur eines RS-Flip-Flop aus zwei NOR-Elementen

Die beschriebene Funktionsweise kann man in übersichtlicher Form im Impulsdigramm darstellen (Bild 5.45). Hier sind die verschiedenen möglichen Zustände an den Ein- und Ausgängen in beliebiger Reihenfolge zeitlich nacheinander angeordnet. Gelangt gleichzeitig an S und R ein Signal mit dem logischen Wert 1, so schalten Q und $\bar{Q}$ in den Nullzustand. Die Zu-



a)

A	B	Q	$\bar{Q}$
0	0	Q_n	$\bar{Q}_n$
1	0	1	0
0	1	0	1
1	1	0	0

b)

Bild 5.45. Zur Funktionsweise eines RS-Flip-Flop

a) Impulsdigramm;
b) Tafel der Übergänge

stände an Q und $\bar{Q}$ sind nicht mehr negiert zueinander. Diese Eingangsbelegung ist deshalb unzulässig (gestricheltes Feld im Bild 5.45). Eine weitere Form der Beschreibung der Eigenschaften des Flip-Flop ist im Bild 5.45b mit der Tafel der Übergänge gegeben. Dabei bedeutet Q_n in der Ausgangsspalte, daß der vorherige Zustand erhalten bleibt. Die Reihenfolge der Zeilen in der Tafel entspricht hier der zeitlichen Reihenfolge im Impulsdiagramm.

Der RS-Flip-Flop kann auch mittels zweier NAND-Elemente aufgebaut werden (Bild 5.46). Aus der Wahrheitstabelle eines NAND folgt, daß bereits bei Belegung eines Eingangs mit dem logischen Wert 0 der Ausgang unabhängig von der Belegung der anderen Eingänge den logischen Wert 1 annimmt. Als Schaltbefehl kommt folglich nur der logische Wert 0 in Frage. Entsprechend der Definition gilt aber für Flip-Flop der logische Wert 1 als Schaltbefehl.

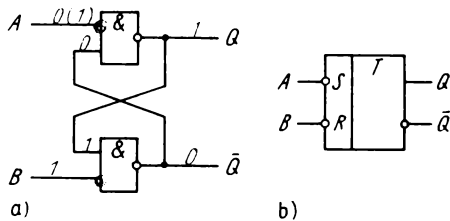


Bild 5.46. RS-Flip-Flop aus zwei NAND-Elementen
a) logische Struktur; b) Schaltungskurzzeichen

Im Bild 5.46b sind deshalb die Befehlseingänge als inverse Eingänge dargestellt. Wird an A oder B ein Signal mit dem logischen Wert 0 angelegt, so erfolgt der Setz- oder Rücksetzvorgang. Für den Speicherbefehl sind die logischen Werte an den einzelnen Ein- und Ausgängen im Bild 5.46b eingetragen. Impulsdiagramm und Tafel der Übergänge (Bild 5.47) sollen auch hier die Vorgänge zusammenfassen.

Nachteil bei den vorstehend beschriebenen Schaltungen ist, daß Störimpulse während der gesamten Betriebszeit ein Stellen der Flip-Flop verursachen können. Sorgt man dafür, daß das Stellen der Flip-Flop nur innerhalb bestimmter Zeitintervalle (Takte) möglich ist, entsteht ein getakteter Flip-Flop. Die logische Struktur eines getakteten RS-Flip-Flop und das zugehörige Schaltungskurzzeichen sind im Bild 5.48 dargestellt. Man erkennt, daß an den Eingängen der Gatter $D3$ und $D4$ nur dann der für die Ausführung des Stellbefehls notwendige Nullzustand vorliegen kann, wenn sowohl am Befehlsein-

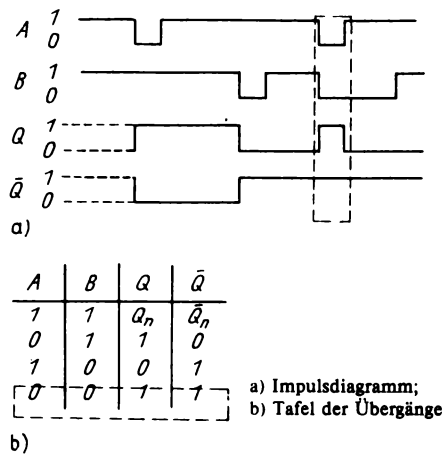


Bild 5.47. Zur Funktionsweise des RS-Flip-Flop nach Bild 5.46.

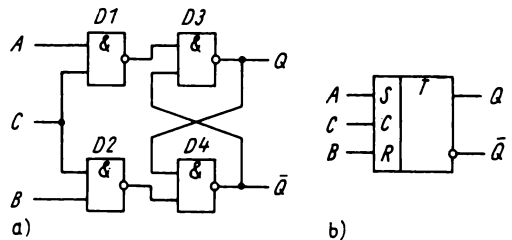


Bild 5.48. Getakteter RS-Flip-Flop
a) logische Struktur; b) Schaltungskurzzeichen

gang wie auch am Takteingang eines der beiden Elemente $D1$ oder $D2$ gleichzeitig Einszustand herrscht. Für Einszustand am Takteingang C kann folglich Einszustand an S Q auf Einszustand setzen. Für Einszustand an C und an R erfolgt das Rücksetzen von Q in den Nullzustand.

5.4.3. JK-Flip-Flop

Im Gegensatz zum RS-Flip-Flop darf beim JK-Flip-Flop gleichzeitig an beiden Speichereingängen der Speicherbefehl liegen. Es erfolgt dabei jeweils ein Umschalten des vorherigen Zustands. Die Funktionsweise eines im Bild 5.49 dargestellten getakteten JK-Flip-Flop soll zunächst durch das Impulsdiagramm

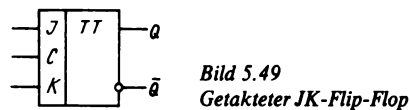


Bild 5.49
Getakteter JK-Flip-Flop

(Bild 5.50) beschrieben werden. Jeweils am Ende des Taktsignals erfolgt die Übernahme des an J oder K liegenden Befehls nach Q und $\bar{Q}$. Befindet sich während der Taktzeit J im Einszustand, so wird beim Übergang des Taktsignals vom Eins- in den Nullzustand an Q der Einszustand entstehen. Bei Einszustand an K kann durch den Takt Q wieder in den Nullzustand gelangen. Haben J und K gleichzeitig Einszustand, so erfolgt bei jeder Taktrückflanke ein Umschalten des Triggerausgangs. Schaltet man ständig J und K auf Einszustand oder J auf Q und K auf $\bar{Q}$ und führt nur den Takteingang heraus, so entsteht ein T-Flip-Flop oder Zähl-Flip-Flop (Bild 5.51). Dieser wirkt wie ein Binärteiler.

In der Tafel der Übergänge (Bild 5.50) ist unter dem Zeitpunkt t_{n+1} der sich nach dem Übergang des Taktsignals aus dem Eins- in den Nullzustand ergebende Ausgangszustand des Flip-Flop eingetragen. Der vor und während des Taktsignals vorliegende Zustand an den Eingängen J und K ist durch den Zeitpunkt t_n gekennzeichnet. Q_n sei der Ausgangswiderstand zum Zeitpunkt t_n . Eine wichtige Ausführungsform des getakteten JK -Flip-Flop ist der Master-Slave-Flip-Flop (master engl. Meister, slave engl. Knecht), dessen logische Struktur im Bild 5.52 dargestellt

ist. Man erkennt, daß zwei getaktete RS -Flip-Flop miteinander verkoppelt sind, wobei der Master durch den Einszustand des Taktes und Slave durch den Nullzustand des Taktes freigegeben werden. Der zeitliche Ablauf der Vorgänge ist im Bild 5.53 dargestellt. Dabei erfolgt im Zeitpunkt 1 das Trennen des Slave vom Master, zwischen den Zeitpunkten 2 und 3 reagiert der Master auf die an J und K liegenden Befehle, wobei wegen der Rückführung vom Slave nur ein Umschalten oder ein Erhalten des vorherigen Zustandes erreichbar ist. Ein mehrmaliges Umschalten des Masters während des Zeitbereiches 2 bis 3 ist nicht möglich. Im Zeitpunkt 4 erfolgt das Übertragen der Information vom Master nach dem Slave und damit an den Ausgang.

Für die erste Zeile der Tafel der Übergänge sind die Zustände an den einzelnen Ein- und Ausgängen im Bild 5.52 eingetragen. Die logischen Werte sind in der zeitlichen Reihenfolge von links nach rechts geordnet. Um den Einfluß von Störimpulsen möglichst klein zu halten, sorgt man für eine möglichst kurze Taktzeit. Die vom Hersteller vorgegebene Mindesttaktzeit darf jedoch nicht unterschritten werden.

Vielfach verfügen JK -Flip-Flop über zusätzliche S - und R -Eingänge mit der im Abschnitt 5.4.2. beschriebenen Wirkung.

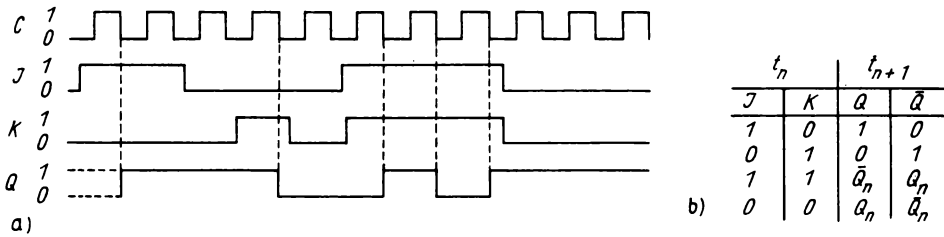


Bild 5.50. Zur Funktionsweise des JK-Flip-Flop
a) Impulsdiagramm; b) Tafel der Übergänge

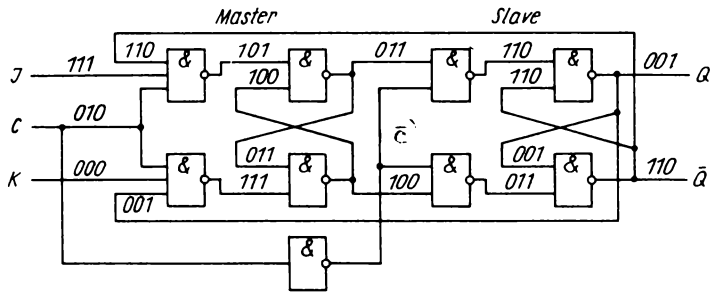


Bild 5.52. Logische Struktur des Master-Slave-Flip-Flop

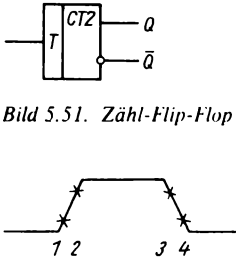


Bild 5.51. Zähl-Flip-Flop
Bild 5.53. Schaltvorgänge in Abhängigkeit vom Taktimpuls

5.4.4.

D-Flip-Flop

Der D-Flip-Flop hat nur einen **Befehlseingang** und einen **Takteingang**. Logische Struktur und Schaltungsskizzen eines statisch gesteuerten D-Flip-Flop zeigt Bild 5.54. Der Unterschied zum getakteten RS-Flip-Flop besteht darin, daß der dortige R-Eingang hier an den Ausgang von Gatter D1 geschaltet wird und sich damit jeweils im zum D-Eingang negierten Zustand befindet. So bedeutet Nullzustand am D-Eingang zwangsläufig Einszustand am mit dem Ausgang von D1 verbundenen Eingang von D2. Es folgt, daß der Flip-Flop immer während des Taktimpulses den gerade an D liegenden Zustand nach Q setzt. Für das Beispiel eines Einszustandes an D sind die jeweiligen Zustände im Bild 5.54 eingetragen.

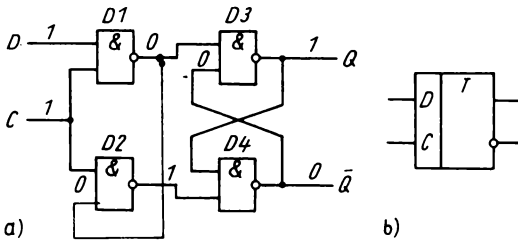


Bild 5.54. Statischer D-Flip-Flop

a) logische Struktur; b) Schaltungsskizzen

Soll eine hohe Sicherheit gegen Störimpulse erreicht werden, verwendet man einen dynamisch gesteuerten D-Flip-Flop (Bild 5.55). Dieser schaltet nur im Zeitpunkt des Übergangs vom

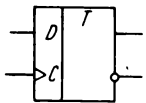
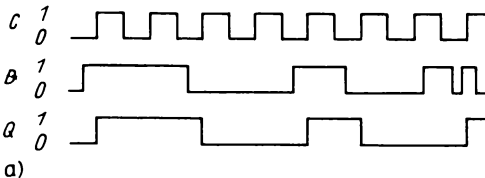


Bild 5.55. Dynamischer D-Flip-Flop



t_n	t_{n+1}	
D	Q	$\bar{Q}$
0	0	1
1	1	0

Bild 5.56. Zur Funktionsweise des dynamischen D-Flip-Flop

a) Impulsdiagramm
b) Tafel der Übergänge

Null- in den Einszustand am Eingang C den zu diesem Zeitpunkt an D liegenden Zustand nach Q. Das Impulsdiagramm und die Tafel der Übergänge sollen auch hier zur Beschreibung dienen (Bild 5.56). Der Zeitpunkt t_n kennzeichnet die Zeit vor und der Zeitpunkt t_{n+1} die Zeit nach dem Übergang des Taktes aus dem Null- in den Einszustand.

5.4.5.

Signalgeneratoren

Signalgeneratoren werden in der Digitaltechnik häufig benötigt. So sind z. B. viele Flip-Flop-Schaltungen getaktet – ein Taktgenerator ist erforderlich. Die Signalgeneratoren sollen je nach Anwendungsfall eine Impulsfolge mit einem bestimmten Tastverhältnis und bestimmter Impulsfolgefrequenz erzeugen. Ein in diskreter Schaltungstechnik ausgeführter Generator wurde bereits im Abschn. 4.3.2. beschrieben. Hier sollen nur mit IS aufgebaute Taktgeneratoren erläutert werden.

Eine einfache Form eines Signalgenerators unter Verwendung zweier TTL-Gatter ist im Bild 5.57 dargestellt. Am Eingang S läßt sich der Generator mit Pegel L stoppen, da dann am Ausgang des Gatters D1 nur noch Pegel H möglich ist. Die Wirkungsweise ist im Impulsdiagramm Bild 5.57 dargestellt. Am Beginn der Erläuterung der Vorgänge liege an y und Punkt A gerade Pegel H. Damit hat Punkt B Pegel L. Über den Widerstand R erfolgt die Umladung des Kondensators (schwarze Strompfeile). Die linke Seite des Kondensators ist dabei zunächst positiver und wird langsam negativer als die rechte Seite. Der erforderliche Umladestrom bewirkt, daß die Spannung an y nach dem Sprung auf Pegel H erst allmählich ihren Höchstwert erreicht. Durch die Umladung wird Punkt A langsam negativer, bis eine Spannung erreicht ist, bei der D1 umgesteuert wird (vgl. Transferkennlinie, Bild 5.16a, die für von H nach L gehenden Eingangspegel bei einer *Schwellschpannung* von etwa 1,4 V einen ausgeprägten Knick aufweist). Die Spannung an B steigt plötzlich an, und Gatter D2 steuert um. Damit entsteht an A ein negativer Spannungssprung, der die auslösende Ursache unterstützt. Der Umkippvorgang erfolgt. Erneut muß der Kondensator umgeladen werden (rote Strompfeile). Der am Anfang hohe Umladestrom bewirkt auch hier, daß an B nicht sofort der Spannungshöchstwert entsteht. Durch die Umladung wird Punkt A wieder positiver,

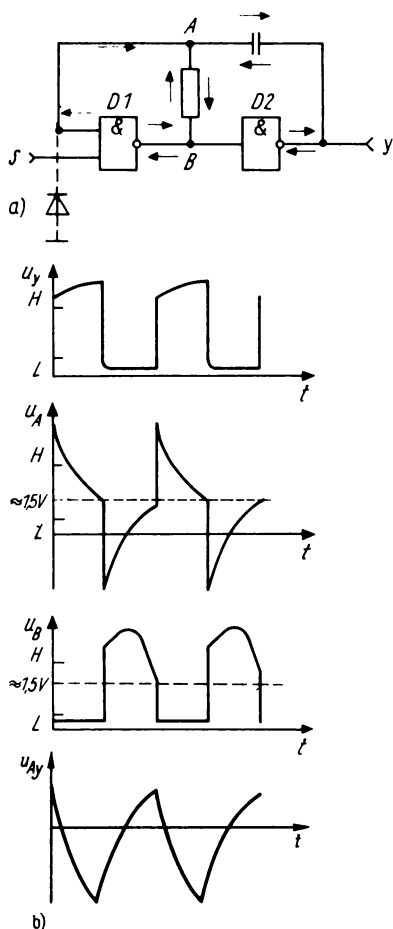


Bild 5.57. Signalgenerator mit zwei TTL-Elementen und einem Kondensator

a) Schaltung; b) Impulsdiagramm bei Verwendung von TTL-Standard

bis Gatter D1 erneut umzusteuern beginnt. Die linke Seite des Kondensators wird jetzt positiver als die rechte. Dabei sinkt gemäß Bild 5.16a die Spannung an B zunächst langsam ab. Ist die Schwellspannung von Gatter D2 erreicht, erfolgt ein erneutes Kippen der Ausgangsspannung. A und y gelangen wieder auf Pegel H. Der Vorgang beginnt von vorn. Mit dieser Schaltung kann nur ein konstantes Tastverhältnis $t_i:T \approx 1:2$ erreicht werden.

Zur sicheren Funktion der Schaltung wird R niedrig gewählt (220Ω). Damit die negative Spannungsspitze an A nicht den Grenzwert für die Gatter überschreitet, sollte man an A eine Diode nach Masse schalten (Bild 5.57a). Die negative Spitze an A wird dann sicher begrenzt. Dabei erfährt auch der HL-Übergang an y eine

geringe Verschlechterung der Flanke. Die Schaltung neigt bei ungünstiger Dimensionierung zur Überlagerung der Ausgangsspannung mit einer hochfrequenten Schwingung.

Beim Einsatz von TTL-LS ergeben sich in der Dimensionierung und im Impulsdiagramm bei sonst prinzipiell gleicher Wirkungsweise einige Änderungen. Diese sind in der nahezu rechteckförmigen Transferkennlinie, den integrierten Clampingdioden sowie der niedrigeren Schwellspannung begründet. So wird der negative Spannungssprung am Punkt A durch die Clampingdioden rasch begrenzt und damit die erforderliche Zeit für die Umladung bis zur Schwellspannung stark verkürzt (siehe Bild 5.58). (Zeitbestimmend ist eine Ladespannungsdifferenz von etwa 1,5 V, während für die Umladung der positiven Spannungsspitze an A ein Spannungsbereich von etwa 3,2 V durchlaufen werden muß.) Das Tastverhältnis wird gegenüber TTL-Standard größer. Für eine sichere Funktion ist R etwa $1 \text{ k}\Omega$ zu wählen.

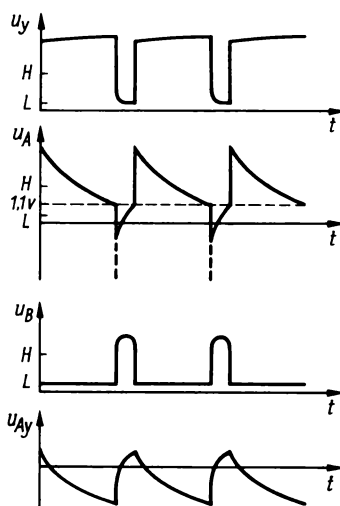


Bild 5.58. Impulsdiagramm zur Schaltung nach Bild 5.57, bei Verwendung von TTL-LS

Für größere Impulszeiten sind größere Kapazitätswerte erforderlich. Wegen der auftretenden Umpolung bei der periodischen Umladung des Kondensators ist kein gepolter Kondensator verwendbar.

Die betrachtete Schaltung ist prinzipiell auch mit CMOS realisierbar. Wegen des wesentlich größeren Ausgangspegelhubes würden hierbei aber an Punkt A die zulässigen Werte weit über-

schritten. Bild 5.59 zeigt eine entsprechende Schaltung mit zusätzlichen Schutzdioden am Eingang von $D1$. $R2$ sollte wesentlich größer als $R1$ gewählt werden.

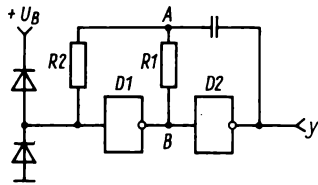


Bild 5.59. Signalgenerator nach Bild 5.57. mit CMOS-Elementen

Eine ähnliche Schaltung lässt sich auch mit MOS-Schaltkreisen realisieren (Bild 5.60). Auch hier erfolgt eine ständige Umladung des Kondensators über R . Dabei darf R nicht zu klein gewählt werden, damit der Ladestrom die Gatterausgänge nicht zu sehr belastet. Wiederum sind die Umladestromverläufe eingetragen. Die Diode am Eingang von Gatter $D1$ ist hier unbedingt erforderlich, damit dieser Eingang im Kippmoment keine positive Spannung annehmen kann. Als zusätzliche Sicherung ist ein Reihenwiderstand eingefügt.

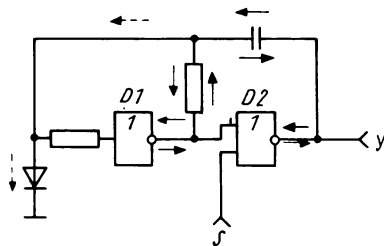


Bild 5.60. Signalgenerator mit zwei MOS-Gattern und einem Kondensator

Eine mit zwei TTL-Gattern und zwei Kondensatoren aufgebaute Schaltung wird im Bild 5.61 gezeigt. Durch unterschiedliche Zeitkonstanten der RC-Glieder ist ein *Tastverhältnis* bis etwa 1:4 möglich. Die Funktion ist im Impulsdiagramm für Pegel H an S dargestellt. Die rot eingetragenen Dioden bleiben zunächst unberücksichtigt. Springt der Pegel an y von H nach L, so springt er an $\bar{y}$ von L auf H. Damit gelangen auch Punkt A auf Pegel H und B auf Pegel L. Der Kondensator $C2$ wird umgeladen. Punkt A wird dabei negativer. Ist die Spannung erreicht, bei der wieder Strom aus dem Eingang von $D1$ zu fließen beginnt, steuert $D1$ um. Dabei wird y

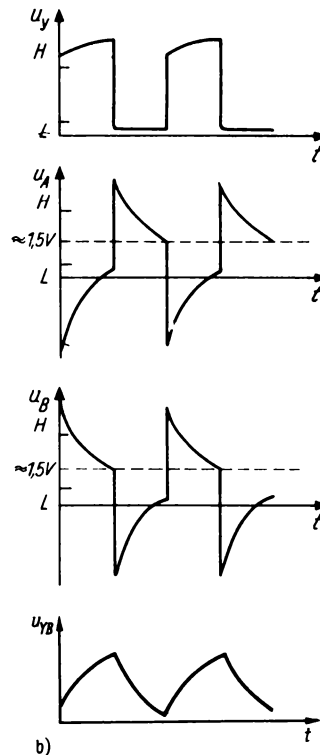
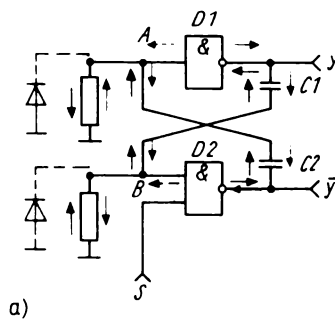


Bild 5.61. Signalgenerator mit zwei TTL-Elementen und zwei Kondensatoren

a) Schaltung; b) Impulsdiagramm

positiver, über Gatter $D2$ wird $\bar{y}$ negativer. Die Wirkung ist der Ursache gleichgerichtet. Der Kippvorgang wird ausgelöst. Gleichzeitig erfolgt eine Entladung von $C1$ und eine Aufladung von $C2$ (schwarze Strompfeile). Durch den steilen Verlauf der Transferkennlinie bei negativer werdendem Eingang und durch die oben beschriebene Mitkopplung erfolgt das Umkippen des Pegels am Ausgang in sehr kurzer Zeit. Es beginnt nun eine Aufladung von $C1$ und eine Entladung von $C2$ (rote Strompfeile). Auch hier wird

der Umkippvorgang ausgelöst, wenn die Spannung am Punkt *B* so weit gesunken ist, daß *D2* umzuschalten beginnt. Der Schwingvorgang kann mit Pegel *L* an *S* unterbrochen werden. $\bar{y}$ nimmt dabei zwangsläufig Pegel *H* an, und *C2* wird stärker als im Schwingbetrieb aufgeladen. Bei erneutem Start ergibt sich zunächst eine längere Periodendauer. Besonders nach Schwingungsunterbrechung wird Punkt *A* kurzzeitig eine negativere Spannung annehmen, als es als Grenzwert zulässig ist. Es ist deshalb zweckmäßig, auch hier an die Eingänge der Gatter Siliziumdioden zu schalten (Bild 5.61a). Ein zu großer Widerstand *R* führt zu einer Verschlechterung des Ausgangsspannungsverlaufs. Der aus den Gattereingängen bei Pegel *L* herausfließende Strom beeinflusst dann zunehmend die Schaltung. Werte von 300 Ω bis 2 k Ω werden für *R* angegeben. Die vorgegebene Schaltung ist bei entsprechender Dimensionierung auch mit anderen Schaltkreisen aufbaubar.

Bei zu niedrigem Wert für *R* schwingt die Schaltung nicht sicher an. Es ist dann möglich, daß nach dem Einschalten an beiden Gatterausgängen Pegel *H* liegt. Ein sicheres Anschwingen garantiert eine Erweiterung der Schaltung gemäß Bild 5.62. Tritt hierbei der Fall ein, daß die Ausgänge der Gatter *D1* und *D2* gleichzeitig auf *H* liegen, so werden über die Gatter *D3* und *D4* die Fußpunkte der Widerstände auf Pegel *H* angehoben. Es wird einer der beiden Gattereingänge eher auf Pegel *H* gelangen – die Ausgangspegel werden zueinander negiert, und die Fußpunkte der Widerstände gelangen wieder auf Pegel *L*. Die Schaltung schwingt so sicher an.

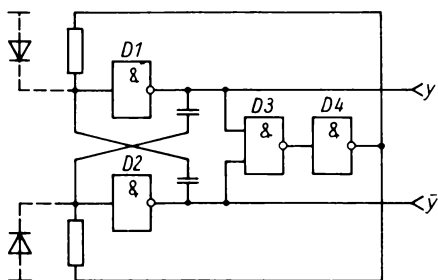


Bild 5.62. Erweiterung der Schaltung nach Bild 5.61. zur Erzielung eines sicheren Anschwingens

Durch Nachschalten eines weiteren Gatters kann man die Flanken der Ausgangsspannung von Generatoren versteilern. Ist das Tastverhältnis

von genau $t_f:T = 1:2$ erwünscht, schaltet man an den Ausgang des Generators einen Flip-Flop (z. B. einen *D*- oder *T*-Flip-Flop). Die Frequenz wird dabei halbiert.

5.4.6.

Monostabile Elemente

Unter einem monostabilen Element versteht man eine Schaltung, bei der durch einen Befehl am Eingang ein Ausgangsimpuls definierter Länge entsteht. Dabei können an den Eingangsbefehl verschiedene Forderungen gestellt sein. So kann verlangt werden, daß der Eingangsbefehl kürzere oder längere Zeit gegenüber der Zeit des Ausgangsimpulses vorliegen muß. Liegt im Befehlseingang ein Differenzierglied, erfolgt je nach Schaltung der Start beim Übergang vom Null- in den Einszustand oder vom Eins- in den Nullzustand am Befehlseingang. Ein einfaches monostabiles Element mit zwei TTL-Gattern ist im Bild 5.63a dargestellt. Das Impulsdiagramm soll die Vorgänge veranschaulichen. Liegt an *x* über längere Zeit Pegel *H*, so liegen an *B* Pegel *L*, an *y* Pegel *H* und somit an *A* Pegel *L* (statischer Zustand). Wird kurzzeitig *x* auf Pegel *L* geschaltet, springen *A* und damit *B* auf Pegel *H*. (An Kondensatoren gibt es keine sprunghaften Spannungsänderungen.) *y* nimmt Pegel *L* an, weshalb auch nach erneutem Übergang von *x* nach Pegel *H* *A* auf Pegel *H* verbleibt. Der Kondensator wird über *R* aufgeladen, bis die Spannung an *B* auf die Schwellspannung von *D2* gesunken ist. Die Spannung an *y* sinkt plötzlich, was über die Rückführungsleitung von *D2* nach *D1* ein Fallen der Spannung an *A* bewirkt und somit die Ursache noch unterstützt. Der Rückkippvorgang erfolgt. Dabei wird *B* negativ, und der Kondensator wird entladen. Erst nach abgeschlossener Entladung ist das monostabile Glied wieder startbereit. Erfolgt schon vor Ablauf dieser *Erholzeit* ein neuer Start, so wird die entstehende *Verweilzeit* kürzer. Die rot eingezeichnete Diode verhindert eine unzulässig hohe negative Spannungsspitze am Eingang von *D2*. Der in den Gatterausgang von *D1* hineinfließende Strom (roter Strompfeil) bewirkt dabei eine Verzögerung des Pegelsprungs von Pegel *H* auf Pegel *L* am Anschluß *A*. Der Aufladewiderstand *R* darf nicht zu groß gewählt werden, damit der bei Pegel *L* aus dem Eingang von Gatter *D2* herausfließende Strom einen Spannungsabfall von höchstens 0,8 V erzeugt. Die geforderten Be-

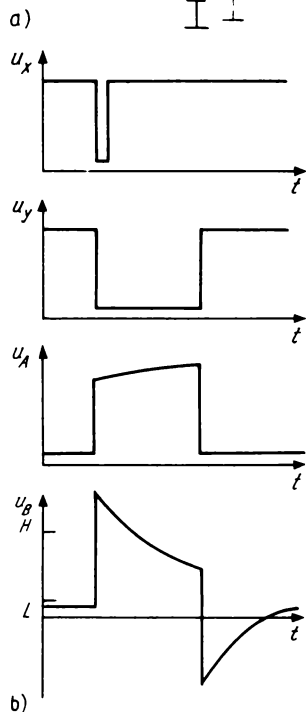
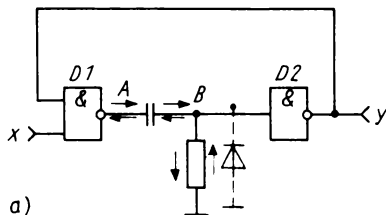


Bild 5.63. Monostabiles Element mit zwei TTL-Elementen

a) Schaltung; b) Impulsdiagramm

triebsbedingungen für den statischen Zustand der Schaltung am Gattereingang von D2 wären sonst nicht mehr gegeben.

Eine Schaltung mit MOS-Schaltkreisen ist im Bild 5.64a dargestellt. Sie benötigt nur kurze Eingangsbeefehle und schaltet beim Übergang von Pegel H nach L am Eingang (Nullzustand in Einzustand am Eingang, da negative Logik). Außerdem ist bei geeigneter Dimensionierung von R die Erholzeit kürzer als die Verweilzeit. Im Impulsdiagramm Bild 5.64b ist am Anfang der statische Zustand dargestellt. Der Kondensator ist aufgeladen und liegt mit seiner linken Seite auf Pegel L und mit seiner rechten Seite auf Pegel H. Springt der Eingang von Pegel H auf Pegel L, so springen A auf Pegel H und y auf Pegel L. Die Rückführung des Pegels L von y auf Gatter D1 sichert, daß auch nach Wegfall

des Eingangssignals an A Pegel H erhalten bleibt. Der Spannungssprung an y verursacht einen entsprechenden an B, und es schließt sich eine Entladung des Kondensators an (schwarze Strompfeile), bis am Eingang des Gatters D2 die Schwellspannung erreicht ist. Damit gelangen Punkt C auf Pegel L und y auf Pegel H, wodurch an Punkt B ein in Richtung der Ursache wirkender positiver Spannungssprung erfolgt. Die Schaltung kippt in ihren Ruhezustand. Damit gelangt Punkt A auf Pegel L, und der Kondensator wird rasch wieder aufgeladen (rote

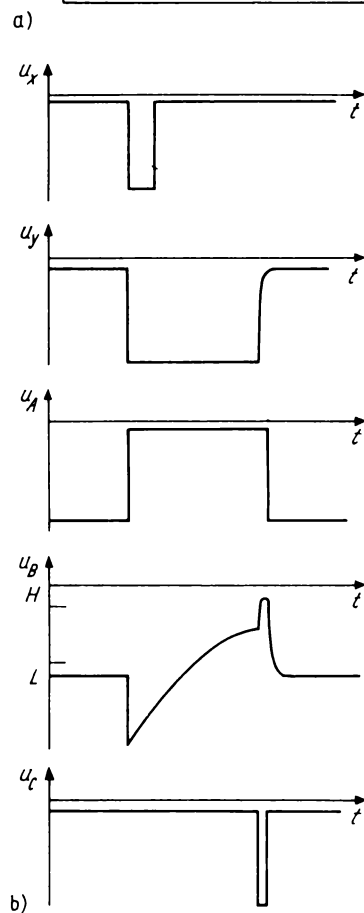
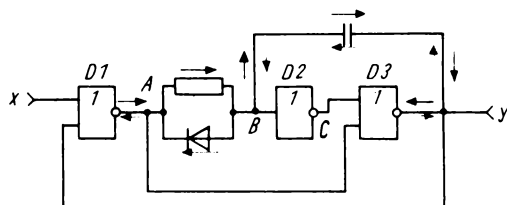


Bild 5.64. Monostabiles Element mit MOS-Schaltkreisen

a) Schaltung; b) Impulsdiagramm

Strompfeile). Punkt *C* nimmt wieder Pegel *H* an. Durch den hohen Ladestrom wird der L-H-Übergang an *y* verzögert. Deshalb kann auch Punkt *B* nicht positiv werden. Eine Überschreitung des zulässigen Grenzwerts kann nicht eintreten.

Eine Variante ohne Einsatz eines die Verweilzeit bestimmenden Kondensators wird im Bild 5.65 gezeigt. Die Verweilzeit wird hier durch die gattereigene Verzögerungszeit festgelegt. Jedes Gatter wechselt den Ausgangspegel nicht gleichzeitig mit, sondern kurze Zeit nach dem Pegelwechsel am Eingang. *Einschalt- und Ausschaltverzögerungszeit* sind meist unterschiedlich. MOS-Schaltkreise in p-Kanal-Hochvolttechnik haben z. B. Verzögerungszeiten von einigen hundert Nanosekunden. Da diese Verzögerungszeiten zu gewollten Zeitverzögerungen bei verschiedenen Anwendungen genutzt wer-

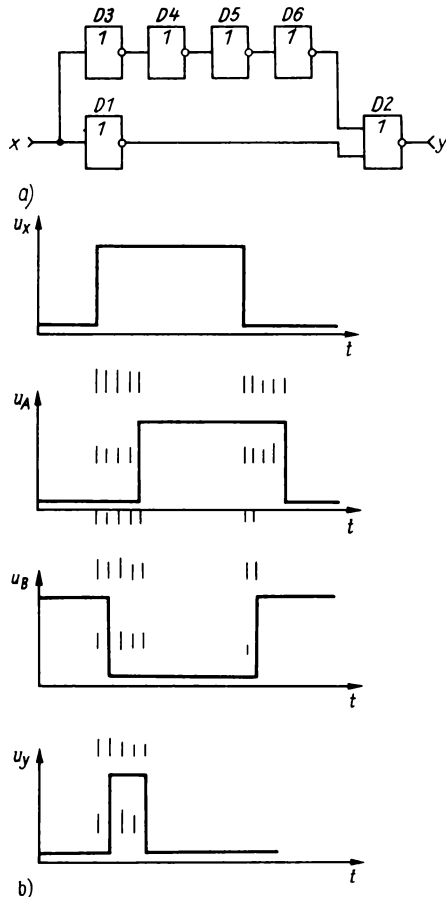


Bild 5.65. Monostabiles Element ohne zeitbestimmende RC-Kombination

a) Schaltung; b) Impulsdigramm

den, wird hier eine entsprechende Schaltung erläutert. Im Ruhezustand soll an *x* Nullzustand herrschen. Über die Gatter *D1* und *D2* hat dann auch *y* Nullzustand, da bereits Einszustand an einem Gattereingang von *D2* an dessen Ausgang Nullzustand bewirkt. Über die Gatter *D3* bis *D6* liegt am zweiten Eingang von *D2* ein Signal mit dem logischen Wert 0. Geht der Eingang *x* vom Nullzustand in den Einszustand über, so wird kurzzeitig an beiden Eingängen von *D2* Nullzustand – an *y* folglich Einszustand – herrschen. Erst nach einer gewissen Verzögerungszeit über die Kette *D3* bis *D6* wird am Ausgang von *D6* Einszustand eintreten: *y* nimmt wieder Nullzustand an. Mit dieser Schaltung sind nur geringe Verweilzeiten erreichbar. Im Impulsdigramm charakterisiert der Abstand zwischen zwei Zeitprojektionslinien die gattereigene Verzögerungszeit.

Es werden auch komplette monostabile Elemente als IS hergestellt. Diese bedürfen dann einer verweilzeitbestimmenden externen Beschaltung mit *C* und *R*. Zur bequemeren Handhabung verfügen diese Schaltkreise meist über mehrere Eingänge mit unterschiedlichen Eigenschaften. So können die die Verweilzeit auslösenden Eingänge dynamisch auf den 0-1-Übergang oder den 1-0-Übergang ansprechen und mit weiteren Eingängen über logische Verknüpfungen verbunden sein (z. B. D 121, DL 123).

5.4.7.

Schwellwertelemente

Bei einem Schwellwertelement soll der Ausgangspegel bei einer bestimmten Eingangsspannung schlagartig von Pegel *L* nach Pegel *H* gehen und bei einer niedrigeren Eingangsspannung wieder von *H* nach *L* zurückkippen. Zwischen Einschaltspannung U_{ein} und Ausschaltspannung U_{aus} besteht eine Differenz, die *Hysteresis* des Schwellwertelements (Bild 5.66).

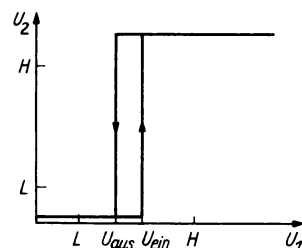


Bild 5.66. Abhängigkeit der Ausgangsspannung eines Schwellwertelements von der Eingangsspannung

Diese Eigenschaft kann man zur Erzeugung einer Rechteckspannung aus einer beliebigen Eingangsschwelspannung nutzen (Bild 5.67). Durch Zusammenschalten eines durchstimmbaren Sinusoszillators mit einem Schwellwertelement kann man eine frequenzdurchstimmbare Rechteckspannung gewinnen. Der Schwellwertabhängige Schaltvorgang wird auch zur Erzeugung einer binären Ausgangsspannung aus einer stetigen Eingangsspannung genutzt. Ergibt z. B. der Meßfühler eines Temperaturregelkreises eine der Temperatur proportionale Spannung, so wird beim Überschreiten einer bestimmten Temperatur (Spannung) das Schwellwertelement umschalten und den Heizkreis abschalten.

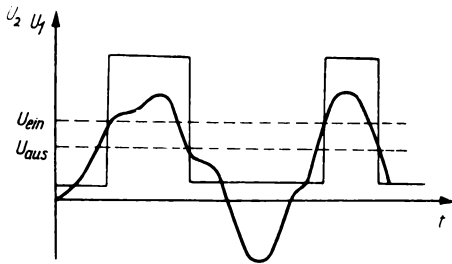


Bild 5.67. Abhängigkeit der Ausgangsspannung eines Schwellwertelements von der Eingangsschwelspannung

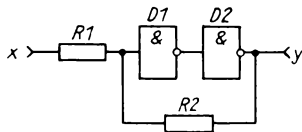


Bild 5.68. Schwellwertelement mit zwei Invertoren

Die Wirkungsweise sei an einer Schaltung mit zwei Invertoren dargelegt (Bild 5.68). Bei einer unterhalb der Schwellspannung von $D1$ liegenden Spannung an x liegt an y Pegel L. Über den Rückführungswiderstand $R2$ wird dieser Zustand noch unterstützt. Steigt jetzt die Spannung an x , so wird trotz Teilung der Eingangsspannung über die beiden Widerstände bei einer bestimmten Eingangsspannung die Schwellspannung von Gatter $D1$ erreicht. Damit sinkt infolge des steilen Bereichs der Transferkennlinie die Spannung am Ausgang von $D1$, was wiederum ein Steigen der Ausgangsspannung des Gatters $D2$ bewirkt. Diese steigende Ausgangsspannung wird über $R2$ auf den Eingang zurückgeführt und unterstützt die auslösende Ursache. Die Schaltung kippt am Ausgang auf Pe-

gel H. Erniedrigt man nun wieder die Eingangsspannung, so wird ein Rückkippen erst bei einer kleineren Spannung erfolgen können, da einer Erniedrigung der Spannung am Eingang des Gatters $D1$ die über $R2$ zurückgekoppelte Spannung entgegenwirkt. Erst wenn infolge weiteren Absinkens der Spannung an x unter die vorherige Einschaltspannung die Spannung am Eingang des Gatters $D1$ wieder die Schwellspannung erreicht, erfolgt der Rückkippvorgang.

Die Schaltschwellen und damit die Hysteresis hängen stark vom Spannungsteiler ab. Bei TTL wählt man $R2$ mit $2,2\text{ k}\Omega$ und erhält in Abhängigkeit von $R1$ die im Bild 5.69 dargestellte Abhängigkeit der Schaltschwellen. Die Einschaltspannung ist fast von $R1$ unabhängig und etwa gleich der Schwellspannung. Bei einem Wert des $R1$ von etwa $800\text{ }\Omega$ wird die Ausschaltspannung negativ, bei $R1 < 100\text{ }\Omega$ ist wegen der dann äußerst klein werdenden Hysteresis mit Instabilität zu rechnen (hierbei erzeugt die Schaltung bei einer Eingangsspannung in Höhe der Schaltschwellen eine wilde Schwingung). Zu beachten ist dabei, daß in den Wert von $R1$ der Innenwiderstand der speisenden Quelle eingeht. Um Fehlfunktionen infolge Störspannungen und Instabilität zu vermeiden, wählt man in der Praxis die Hysteresis möglichst groß.

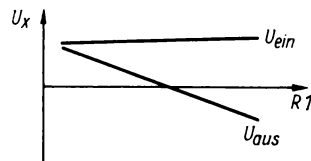


Bild 5.69. Schaltschwellen der Schaltung nach Bild 5.68. in Abhängigkeit von $R1$

Die Eingangsspannung am Gatter $D1$ darf die zulässigen Grenzwerte nicht überschreiten. Gegen eine negative Spannung schützt eine vom Eingang von $D1$ nach Masse geschaltete Diode. Bei TTL-Standard ist zusätzlich zu verhindern, daß die Eingangsspannung von $D1$ positiver als $5,5\text{ V}$ wird.

Schwellwertelemente können auch innerhalb eines integrierten Schaltkreises integriert sein. Als Beispiel hierfür sei die entsprechende Teilschaltung des Schaltkreises A 301 im Bild 5.70 dargestellt. $V9$ und $V10$ stellen eine Differenzverstärkerstufe dar, in die die mit $V11$ gebildete Konstantstromquelle einspeist. Ist die Spannung an x gleich der innerhalb des A 301 stabi-

lisierten Spannung von 2,9 V, so ist V_{12} aufgesteuert. Damit ist der Spannungsabfall an R_{13} groß, und V_{13} ist ebenfalls voll durchsteuert. An y liegt nur U_{CEsat} von V_{13} . Der Transistor V_{10} ist durch den großen Spannungsabfall an R_{15} gesperrt. Sinkt die Spannung an x , so wird ab einer bestimmten Spannung der Differenzverstärker aus der Übersteuerung in den aktiven Bereich gelangen und der Strom durch V_9 und V_{12} sinken. Das bewirkt, daß auch der Strom durch V_{13} kleiner wird und daß das Potential an der Basis von V_{10} steigt. Der somit steigende Strom durch V_{10} wirkt in Richtung der Ursache (sinkender Strom durch V_9); es kommt zum Sprung der Ausgangsspannung auf etwa 2,9 V. Erhöht man nun wieder die Spannung an x , so wird der Strom durch V_9 wieder zu steigen beginnen, was über V_{12} eine Erhöhung des Stromes durch V_{13} und damit eine Erniedrigung des Stromes durch V_{10} bewirkt. Die Spannung am Ausgang springt wieder auf etwa 0,2 V.

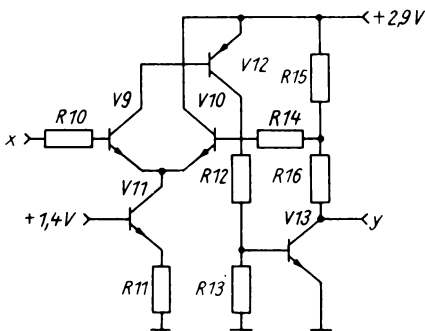


Bild 5.70. Teilschaltung „Schwellwertelement“ des integrierten Schaltkreises A 301

Schwellwertelemente werden auch als selbständige IS gefertigt. Bei diesen Schaltkreisen ist zu unterscheiden, ob sich die Schaltspannungen proportional mit der Betriebsspannung ändern oder ob sie betriebsspannungsunabhängig konstant sind. Je nach Ausführungsform ist die Hysterese beeinflussbar.

Zusammenfassung

Zu den sequentiellen Schaltungen gehören die Flip-Flop. Sie sind je nach Herstellung der logischen Beziehungen eingeteilt. Die Ausgänge Q und $\bar{Q}$ sind zueinander negiert und ändern ihre Zustände nur als Folge von Stellbefehlen an den Eingängen. Die jeweiligen Eigenschaften werden zweckmäßig durch das Impulsdiagramm oder die Tafel der Übergänge beschrieben. Die

Anwendung eines einzelnen Flip-Flop kann als Speicher oder als Binärteiler erfolgen.

Die Signalgeneratoren können eine Rechteckspannung mit bestimmtem Tastverhältnis $t_f:T$ erzeugen. Ihre Wirkungsweise beruht auf der Umladung von Kondensatoren und einer extrem hohen Mitkopplung. Die Gefahr der möglichen Überschreitung der Grenzwerte für die Eingangsspannung der Gatter ist zu beachten; eventuell sind besondere schaltungstechnische Maßnahmen erforderlich.

Ein monostabiles Element erzeugt nach einem entsprechenden Eingangsbefehl einen Ausgangsimpuls definierter Länge. Hierbei ist besonders der Zusammenhang zwischen Verweilzeit und Erholzeit zu beachten. Die Zeitdauer des Ausgangsimpulses kann wiederum durch die Umladung eines Kondensators oder durch Nutzung der gattereigenen Einschalt- und Ausschaltverzögerungen bestimmt werden. Der Kippvorgang ist auch hier die Folge einer starken Mitkopplung.

Bei einem Schwellwertelement kippt der Ausgangspegel jeweils bei einer bestimmten Eingangsspannung. Zwischen Einschalt- und Ausschaltspannung besteht eine bestimmte Hysterese, die aus Stabilitätsgründen nicht zu klein gewählt sein sollte. Der Einsatz kann u. a. zur Gewinnung einer binären Ausgangsspannung aus einer stetigen Eingangsspannung, zur Herstellung einer Rechteckspannung aus einer beliebigen Wechsellspannung sowie zur Signalregenerierung erfolgen. Auch hier sorgt eine Mitkopplung für das Kippen der Ausgangsspannung.

5.5.

Ausgewählte Teilschaltungen digitaler Geräte

5.5.1.

Signaleingabe über mechanische Kontakte

Bei manueller Signaleingabe verwendet man heute noch häufig Taster (z. B. Mikroschalter). Beim Betätigen des Schaltkontakts erfolgt nicht nur ein einmaliges sofortiges Schalten, sondern es treten innerhalb einer sehr kurzen Zeit mehrere Schaltschritte auf. Diese Erscheinung bezeichnet man als *Kontaktprellen*. Ohne besondere Schaltmaßnahmen würde das Kontaktprellen zu Fehlschaltungen führen. So würde ein

Zähler nicht nur den gewünschten Zählimpuls, sondern die Gesamtzahl der Prellungen zählen. Eine einfache Entprellschaltung ist mit einem RS-Flip-Flop aufgebaut (Bild 5.71). Durch den Umschalter wird jeweils ein Eingang des Flip-Flop auf Nullzustand gebracht; der andere Eingang gelangt über den Widerstand in Einszustand. Sollte am Kontakt ein Prellen auftreten, so wird dies am Flip-Flop-Ausgang nicht wirksam, da der Flip-Flop bereits beim ersten Impuls umschaltet. Auch ein Prellen beim Öffnen des vorher geschlossenen Kontakts bleibt wirkungslos, da der Flip-Flop erst beim Schließen des anderen Kontakts umgeschaltet werden kann.

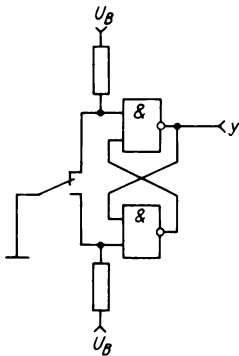
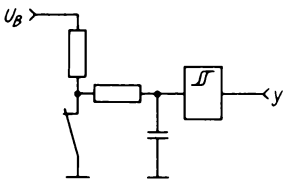
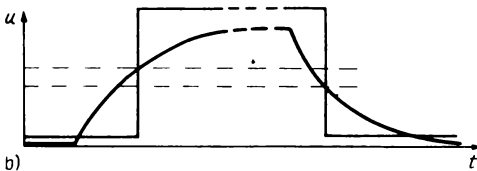


Bild 5.71. Entprellschaltung mit RS-Trigger

Eine weitere Schaltung nutzt die Wirkung eines Integrierglieds aus (Bild 5.72). Durch den Kondensator wird eine Spannungsänderung am Eingang des Schwellwertelements so verzögert, daß das Kontaktprellen nicht wirksam wird. Die Zeitkonstante sollte etwa 5 ms betragen. Da durch das RC-Glied die Anstiegs- und Abfallzeit des Signals stark verzögert werden, ist ein



a)



b)

Bild 5.72. Entprellschaltung mit Schwellwertelement

a) Schaltung; b) Impulsdiagramm

Schwellwertelement zur Signalflankenregenerierung nachgeschaltet. Im Impulsdiagramm sind die Spannung am Ausgang des Integrierglieds schwarz und die Spannung am Ausgang des Schwellwertelements rot dargestellt. Bei dieser Schaltung tritt eine Signalverzögerung auf, die jedoch bei manueller Signaleingabe ohne Bedeutung ist.

Ein monostabiles Element kann man ebenfalls benutzen, um die Auswirkung des Kontaktprellens auf die Schaltung zu verhindern. Die Verweilzeit muß dabei größer als die maximale Prelldauer sein.

5.5.2.

Impulsverzögerung und Flankenversteilerung

Eine Impulsverzögerung kann notwendig sein, wenn eine Schaltung erst nach Ablauf eines Einschwingvorgangs gestellt werden soll oder wenn ungleiche Signallaufzeiten mehrerer gleichzeitig zu verarbeitender Signale auszugleichen sind. Die Verzögerung durch Kettenschaltung von Gattern ist besonders zum Ausgleich unterschiedlicher Signallaufzeiten geeignet (s. auch Bild 5.65). Bei gewünschter größerer Verzögerungszeit wird der Aufwand mit der Kettenschaltung zu groß. Hier ist eine Schaltung nach Bild 5.73 einsetzbar.

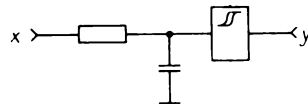
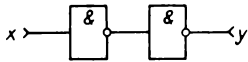
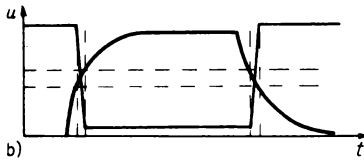


Bild 5.73. Impulsverzögerung durch ein Integrierglied

Durch Parallelkapazitäten kommt es zu unerwünschten Flankenverschleifungen. Das gilt besonders für solche Schaltungen, bei denen der Ausgang kapazitiv belastet wird (s. Abschnitte 5.4.5. und 5.4.6.). Eine Flankenversteilerung kann man mit einem Schwellwertelement erreichen (s. Bilder 5.72 und 5.73). Die steile Transferkennlinie eines Inverters gestattet eine Lösung gemäß Bild 5.74. Im zugehörigen Impulsdiagramm ist die Eingangsspannung schwarz dargestellt. Die markierten Spannungswerte entsprechen dem Spannungsbereich am Eingang eines Inverters, in dem die Ausgangsspannung den Pegelwechsel durchläuft. Rot dargestellt ist die Ausgangsspannung nach dem ersten Inverter. Bei dieser Schaltung sollte die Anstiegs- und Abfallzeit der Eingangsspannung nicht zu



a)



b)

Bild 5.74. Flankenversteilerung durch zwei Inverter

a) Schaltung; b) Impulsdiagramm

groß sein, da sonst Instabilität möglich ist. Für TTL-Schaltkreise nach Abschn. 5.3.2. soll z. B. die Anstiegs- und Abfallzeit für kombinatorisch wirkende Schaltkreise kleiner als $1 \mu\text{s}$ sein.

5.5.3.

Anschluß systemfremder Lasten

An der Peripherie einer digitalen Schaltung müssen die gewonnenen Ausgangssignale häufig entweder angezeigt werden (z. B. über Lampen) oder Schaltungen der Leistungselektrik ansteuern (z. B. Leistungsthyristoren oder Relais). Zur Anpassung an die externen Einheiten sind wiederum Interface-Schaltungen notwendig (s. Abschn. 5.3.3.), da alle externen Einheiten systemfremde Lasten darstellen. Am Übergang zur Peripherie finden sich heute bevorzugt TTL-Schaltkreise oder Schaltkreise anderer Schaltkreisfamilien mit TTL-kompatiblen Ausgang. Deshalb sollen hier nur Schaltungen dargestellt werden, bei denen die systemfremde Last durch TTL-Schaltkreise angesteuert wird.

Ein TTL-Gatter mit „totem pole output“ (s. Bild 5.9) kann bei Pegel H am Ausgang Strom abgeben (engl., pull up) oder bei Pegel L Strom aufnehmen (engl., pull down). Die Stromaufnahme-fähigkeit ist bei $N_0 = 10$ mit 16 mA größer als die Stromabgabefähigkeit mit $0,4 \text{ mA}$. Deshalb sollte die Ansteuerung systemfremder Lasten möglichst durch Nutzung der Stromaufnahme-fähigkeit erfolgen. Als weiterer Nachteil tritt bei Nutzung der Stromabgabe innerhalb des steuernden Schaltkreises eine erhebliche Verlustleistung am vom Laststrom durchflossenen Widerstand R_4 und an V_5 (Bild 5.75) auf. Die Wahl der Bezeichnungen erfolgte in Anlehnung an Bild 5.8.

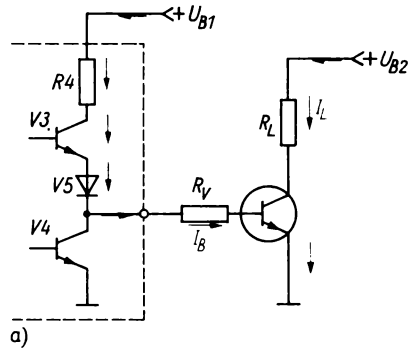
Ein Beispiel mit Nutzung der Stromabgabefähigkeit wird im Bild 5.75 gezeigt. Der Transistor

wirkt hierbei als Stromverstärker. R_V wird berechnet nach

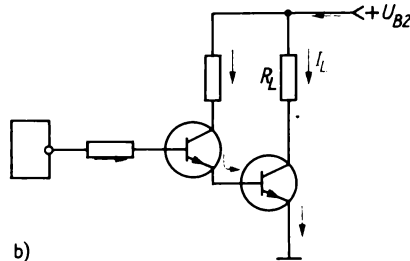
$$R_V = \frac{U_{OH} - U_F}{I_B} \approx \frac{2,4 \text{ V} - 0,7 \text{ V}}{I_B}$$

$$R_V \approx \frac{1,7 \text{ V}}{I_B} \quad (5.4)$$

Dabei stellt I_B den für I_L erforderlichen Basisstrom dar. I_B sollte aus Sicherheitsgründen stets etwas überdimensioniert sein. (Man beachte, daß die Stromverstärkung eines Transistors nicht nur vom verwendeten Exemplar – gekennzeichnet durch den Stromverstärkungsbereich innerhalb der Stromverstärkungsgruppe –, sondern auch vom erforderlichen Kollektorstrom und von der Kristalltemperatur abhängt.) Reicht die Stromabgabefähigkeit des Schaltkreises für I_B nicht aus, wird eine Darlingtonschaltung nach Bild 5.75 eingesetzt.



a)



b)

Bild 5.75. Ansteuerung einer systemfremden Last durch Stromabgabe

a) mit Transistor; b) mit Darlington-Schaltung

Man kann, wenn ausschließlich die systemfremde Last angeschlossen wird und die Einhaltung des Pegels nicht garantiert werden muß, auch höhere Ströme entnehmen. (Vom Hersteller wird für TTL-LS ein für diesen Fall höchstzulässiger Strom von 8 mA angegeben.) Auch

der Einsatz eines *pull-up-Widerstandes* zur Erhöhung der Stromabgabefähigkeit ist möglich ($R1$ auf Bild 5.76). Der Wert dieses Widerstandes ist mindestens so hoch zu wählen, daß bei Pegel L am Ausgang der in den Schaltkreisausgang hinein fließende Strom den zulässigen Wert nicht überschreitet ($R > 360 \Omega$). Unter den genannten einschränkenden Bedingungen soll $R2 > 120 \Omega$ gewählt werden. Für TTL-LS ergibt sich dann ein möglicher Basisstrom von 15 mA. Eine Schaltung mit Nutzung der Stromaufnahmefähigkeit wird im Bild 5.77 gezeigt. Der Wi-

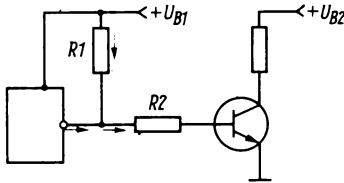


Bild 5.76. Erhöhung der Stromabgabefähigkeit durch pull-up-Widerstand

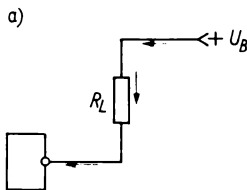
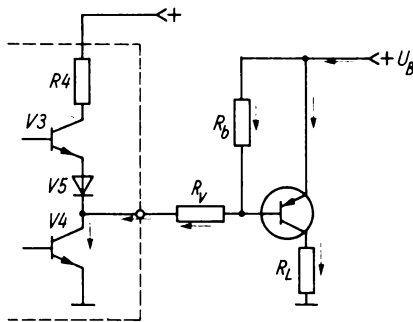


Bild 5.77. Ansteuerung einer systemfremden Last durch Stromaufnahme

a) mit einem Transistor; b) direkt

derstand R_b leitet den bei gesperrtem Transistor $V4$ noch fließenden Reststrom ab. Sie ist auch bei Ausgängen mit offenem Kollektor (engl., open collector output) einsetzbar. Reicht die maximale Stromaufnahmefähigkeit zum Betreiben der Last aus (z. B. bei Ansteuerung einer LED), so kann man auf einen Verstärkertransistor verzichten (Bild 5.77b). Bei Schaltkreisen mit offenem Kollektor sind entsprechend der maximalen Spannungsfestigkeit des Ausgangstransistors auch Lasten mit höherer Betriebsspannung schaltbar (Bild 5.78). Bei Ausgangspegel L fließt der schwarz eingezeichnete Strom; an der Basis des Verstärkertransistors liegt nur U_{CEsat} des Ausgangstransistors, und der Verstärkertransistor ist gesperrt. Bei Ausgangspegel H fließt der rot eingezeichnete Strom.

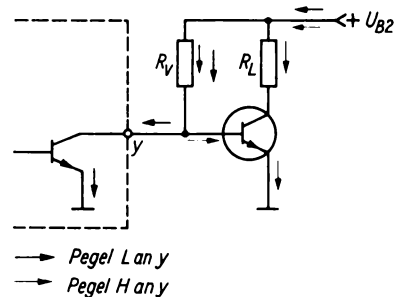


Bild 5.78. Ansteuerung einer systemfremden Last durch einen Ausgang mit offenem Kollektor

5.5.4. Dual- und Dezimalzähler

Man unterscheidet *synchrone* und *asynchrone* Zähler. Beim *asynchronen Zähler* wird der Zählimpuls nur an den ersten Flip-Flop geführt. Alle weiteren Flip-Flop erhalten ihre Zählimpulse vom Ausgang des vorhergehenden Flip-Flop. Damit schalten die aufeinanderfolgenden Stufen jeweils um eine Schaltzeit verzögert (asynchron). Ein Asynchrone Zähler mit vier Stufen ist im Bild 5.79 dargestellt. Alle J- und K-Eingänge haben Einzustand. Es sind 2^n verschiede-

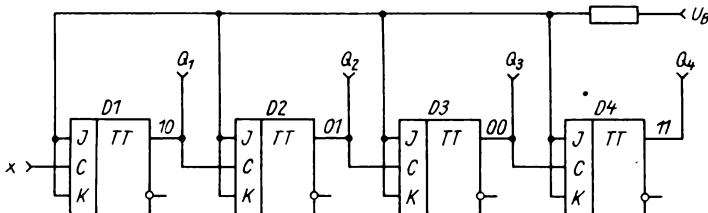


Bild 5.79. Asynchrone Zähler mit JK-Flip-Flop, dual

dene Ausgangszustände möglich (n Anzahl der Flip-Flop). Die dargestellte Schaltung ermöglicht demnach 16 verschiedene Zählerstände, damit ein Zählen von 0 bis 15. Sind alle möglichen Zustände genutzt, so spricht man von einem *Dualzähler*. Zum besseren Verständnis ist im Bild 5.80 das Impulsdigramm eines Dualzählers angegeben und die Übergänge vom Zählerschritt 4 nach 5 und 9 nach 10 sind besonders markiert. Für den Zählerschritt 9 nach 10 sind im Bild 5.79 die Zustände an den Ausgängen $Q1$ bis $Q4$ angegeben. Man erkennt, daß an den Takteingängen C von $D1$ und $D2$ ein 1-0-Übergang liegt; diese Flip-Flop schalten am Ausgang folglich um. Die Flip-Flop $D3$ und $D4$ können nicht umschalten, da an ihren Takteingängen kein 1-0-Übergang liegt. In gleicher Weise sind alle weiteren Zählschritte überprüfbar.

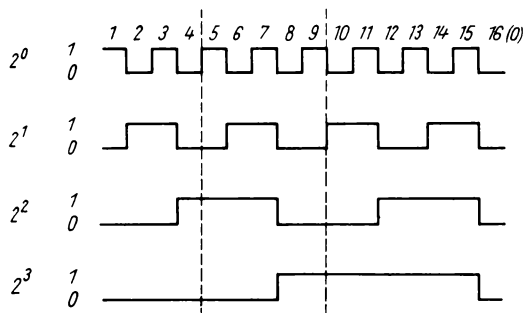


Bild 5.80. Impulsdigramm eines vierstufigen Dualzählers

Bei einer Schaltung nach Bild 5.81 wird der Zähler schon nach dem neunten Schritt auf 0 gestellt (*Dezimalzähler*). An den Ausgängen $Q1$ bis $Q4$ sind wiederum für den Zählschritt 9 nach 0 die logischen Werte angegeben. Der Vergleich zwischen dem Zählerstand 10 des Dualzählers und 0 des Dezimalzählers in Bild 5.82 zeigt die notwendigen Änderungen beim Dezimalzähler gegenüber dem Dualzähler. So darf $Q2$ beim Zählerstand 10 nicht auf 1 gesetzt werden,

den, muß aber sowohl beim Zählerstand 2 als auch 6 auf 1 gesetzt werden können. Es ist dazu der Eingang J von $D2$ so zu beschalten, daß für die Zählerstände 1 und 5 der logische Wert 1 anliegt, beim Zählerstand 9 hingegen der logische Wert 0. Diese Bedingung erfüllt die Rückführung des Ausgangs $\bar{Q}4$ (0...7 logischer Wert 1, ab 8 logischer Wert 0). Als weitere notwendige Änderung gegenüber dem Dualzähler er-

Ausgang	Wertigkeit	Zustand bei Zählerstand		
		9	10	0
$Q1$	2^0	1	0	0
$Q2$	2^1	0	1	0
$Q3$	2^2	0	0	0
$Q4$	2^3	1	1	0

Bild 5.82. Tafel der logischen Werte zum Dezimalzähler

kennt man, daß $Q4$ bereits beim Zählerstand 10 auf 0 zurückschalten muß. Beim Übergang zum Zählerschritt 10 (0) liegt aber nur am Ausgang von $D1$ der zum Schalten erforderliche 1-0-Übergang, weswegen der Takteingang von $D4$ mit Ausgang $Q1$ zu verbinden ist. Um nun ein unerwünschtes Schalten von $D4$ nach den Zählerständen 1, 3 und 5 zu verhindern, werden die J -Eingänge von $D4$ mit den Ausgängen $Q2$ und $Q3$ verbunden, die nur ab Zählerstand 6 gemeinsam den logischen Wert 1 führen.

Auch mit D -Flip-Flop sind asynchrone Zähler leicht aufbaubar. Ein Beispiel wird im Bild 5.83 gezeigt. Die D -Eingänge sind jeweils mit $\bar{Q}$ des gleichen Flip-Flop verbunden. Es kann Q nur dann schalten, wenn am Takteingang ein Übergang vom Null- in den Einszustand vorliegt. Zur Erläuterung sei der Übergang vom Zählerstand 4 zum Zählerstand 5 näher betrachtet. Bei Zählerstand 4 hat $Q3$ Einszustand, $Q1$, $Q2$ und $Q4$ haben Nullzustand. Beim folgenden

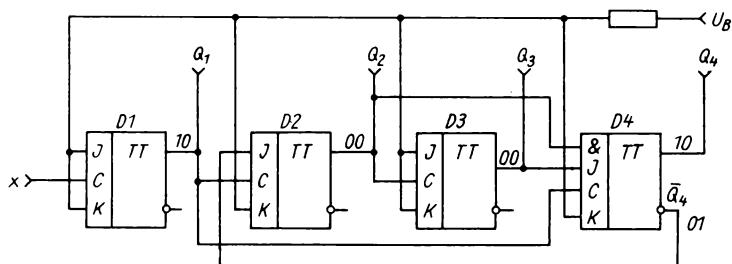


Bild 5.81. Asynchronzähler mit JK-Flip-Flop, dezimal

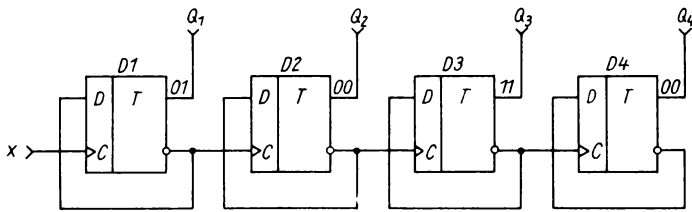


Bild 5.83. Asynchronzähler mit D-Flip-Flop, dual

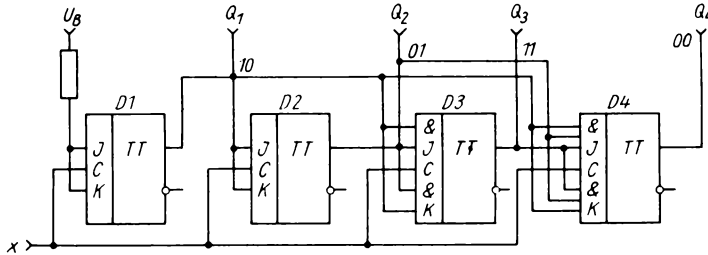


Bild 5.84. Synchronzähler mit JK-Flip-Flop, dual

0-1-Übergang am Zählereingang kann $Q1$ auf Einszustand schalten; $Q2$ behält den vorherigen Zustand, da an seinem Takteingang ein 1-0-Übergang liegt; $Q3$ und $Q4$ behalten ebenfalls den vorherigen Zustand, da an ihren zugehörigen Takteingängen keine Signalübergänge vorliegen.

Bei *Synchronzählern* wird der Zählimpuls dem Takteingang aller Flip-Flop parallel zugeführt. Das Umschalten der Flip-Flop erfolgt gleichzeitig (synchron). Damit die Flip-Flop entsprechend der Wertigkeit ihrer Ausgänge und nicht etwa gleichzeitig schalten, sind entsprechende Leitungen zu den J- und K-Eingängen zu führen. Eine Möglichkeit ist im Bild 5.84 dargestellt. Zur Erklärung sei der Zählschritt von 5 nach 6 gemäß Bild 5.80 betrachtet. Nach dem fünften Zählschritt hat $Q1$ Einszustand; damit kann $Q2$ beim nächsten Zählimpuls auf Einszustand gesetzt werden. Da an $Q2$ Nullzustand herrscht, sind über die J- und K-Eingänge die Flip-Flop $D3$ und $D4$ so verriegelt, daß sie nicht umschalten können. Beim Zählschritt von 7 nach 8 können die Ausgänge $Q2$, $Q3$ und $Q4$ umschalten, da an ihren J- und K-Eingängen entsprechend der Schaltung Einszustand vorliegt.

In der Praxis werden Zählerschaltungen häufig benötigt. Die Bauelementehersteller bieten deshalb komplette Zähler als integrierte Schaltkreise an. Diese Zähler verfügen zur bequemeren Nutzung über weitere Ein- und Ausgänge als die oben beschriebenen Grundsaltungen. Sollen Dezimalzähler zu einem mehrstelligen

Zähler verbunden werden, so ist zur Ansteuerung des jeweils höherwertigen Digits ein *Übertragsimpuls* vom Ausgang des niederwertigeren Digits notwendig (engl., *carry output*). Soll der Zähler sowohl vorwärts als auch rückwärts zählen können, muß er zusätzlich über einen Ausgang für den Übertrag rückwärts verfügen (engl., *borrow output*).

Ebenso besteht häufig die Forderung, daß zu Beginn der Zählung der Zählerstand 0 oder ein beliebig wählbarer Zustand einstellbar sein soll. Dazu dienen Rücksetzeingang (engl., *reset*) oder Ladeeingang (engl., *load, set*) sowie Dateneingänge.

Da somit der Ausgangszustand des Zählers von verschiedenen Eingängen beeinflusst werden kann, ist bei der praktischen Anwendung die *Priorität* der Eingänge zu berücksichtigen. Voreinstelleingänge haben allgemein die höhere Priorität als die Zählereingänge, bei den verschiedenen Voreinstelleingängen hingegen ist von Schaltkreis zu Schaltkreis die Priorität unterschiedlich.

Abschließend soll ein Beispiel eines konkreten Zählerschaltkreises betrachtet werden. Der Schaltkreis D 192 (DL 192) stellt einen synchronen vor-/rückwärts Dezimalzähler dar (Bild 5.85). Im Impulsdiagramm erkennt man, daß am Übertragsausgang vorwärts (CR1) beim 1-0-Übergang des Zähltrittes 9 ein 1-0-Übergang entsteht und beim 0-1-Übergang des Zähltrittes 0 (10) der Übertragsausgang wieder den Wert 1 annimmt. Mit diesem Übergang kann das höherwertige Digit geschaltet werden.

Neben der Möglichkeit, den Zähler über reset auf 0 voreinstellen zu können, ist das Laden eines beliebigen, an den Dateneingängen anliegenden Zählerstandes durch den logischen Wert 0 am Seteingang möglich. (Zu beachten ist, daß bei offenem Eingang *R* dieser auf Pegel H liegt, damit den logischen Wert 1 annimmt und ein ständiges Rücksetzen des Zählers auf 0 erfolgt. Der Eingang *R* ist folglich, wenn er nicht benötigt wird, auf Masse zu schalten).

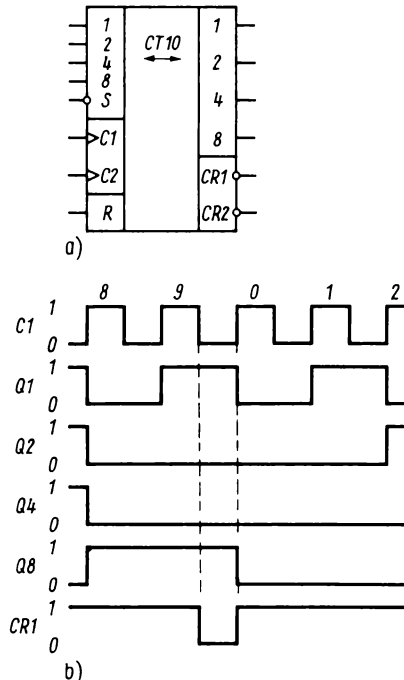


Bild 5.85. Dezimalzählerschaltkreis D 192 (DL 192)

a) Schaltungskurzzeichen; b) Impulsdiagramm

5.5.5.

Frequenzteiler

Bei der Frequenzaufbereitung verwendet man vielfach Frequenzteiler, mit deren Hilfe man, von einer Mutterfrequenz ausgehend, niedrigere Frequenzen mit ganzzahligen Teilverhältnissen gewinnen kann. Zu unterscheiden sind auch hier asynchrone und synchrone Frequenzteiler. Betrachtet man den Ausgang *Q4* des Dualzählers nach Bild 5.79, so erkennt man, daß an *Q4* nur noch $\frac{1}{16}$ der am Zähleringang liegenden Frequenz vorliegt. Mit dem Dezimalzähler nach Bild 5.81 wäre eine Frequenzteilung 10:1 möglich. Bei der dort beschriebenen

Schaltung gelangt *Q4* beim achten Eingangsimpuls in den Einszustand, um bereits beim zehnten (nullten) Eingangsimpuls wieder Nullzustand anzunehmen. Die Ausgangsspannung an *Q4* hat folglich kein Tastverhältnis von $t:T = 1:2$ (symmetrische Rechteckspannung), sondern ein Tastverhältnis von 1:5. Die sechs überflüssigen Zähler Schritte der Schaltung nach Bild 5.81 werden durch Überspringen der Ausgangszustände 10 bis 15 unwirksam. Es ist durch andere Rückführungsleitungen möglich, die nicht benötigten Zähler Schritte an beliebiger Stelle des Zählvorgangs zu überspringen; dabei entsteht jeweils ein anderes Tastverhältnis der Ausgangsspannung.

Ein synchroner Frequenzteiler 16:1 ist mit einer Anordnung nach Bild 5.84 realisierbar. Ein Beispiel für einen synchronen Frequenzteiler 5:1 ist im Bild 5.86 dargestellt.

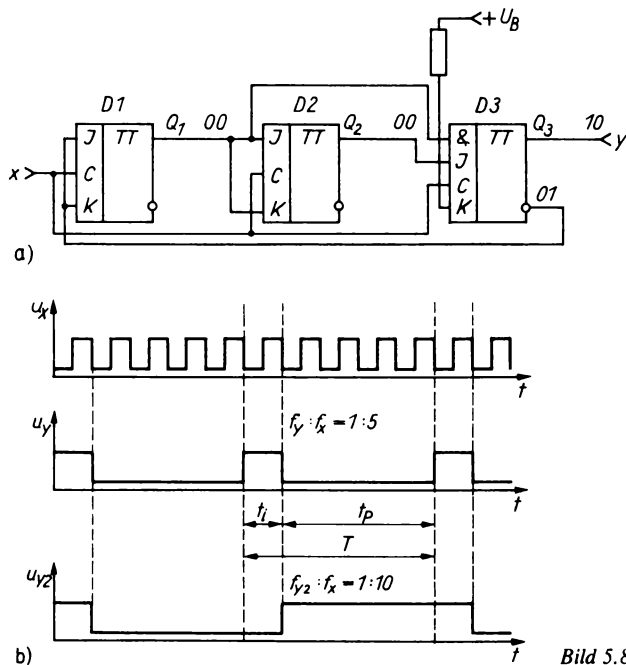
Ohne Überspringen von Zähler Schritten würde bei einem dreistufigen Zähler eine Frequenzteilung von 8:1 entstehen. Zur Reduzierung des Teilverhältnisses sollen im Beispiel die Zähler Schritte 5 bis 7 übersprungen werden. Die dazu notwendigen Änderungen ergeben sich aus Bild 5.86c. Beim Übergang vom Zählerstand 4 zum nächstfolgenden darf *Q1* nicht auf 1 schalten und muß *Q3* auf 0 schalten. Ersteres erreicht man, indem man *J* und *K* von *D1* mit $\bar{Q3}$ verbindet, wo ab Zählerstand 4 der logische Wert 0 liegt. Das veränderte Schaltverhalten an *Q3* wird über die beiden *J*-Eingänge erreicht, die beim betrachteten Übergang beide den logischen Wert 0 haben. Frühestens nach dem Zählerstand 3 kann hingegen *Q3* auf 1 schalten, da erst ab Zählerstand 3 *Q1* und *Q2* den logischen Wert 1 haben. Die Spannung an *Q3* weist dabei ein Tastverhältnis $t:T = 1:5$ auf.

Ist ein Tastverhältnis der geteilten Frequenz von 1:2 gefordert, kann man dem Teiler einen Zähl-Flip-Flop nach Bild 5.51 nachschalten. Die dann entstehende Ausgangsspannung ist im Bild 5.86b mit u_{y2} bezeichnet. Dabei verringert sich jedoch die Ausgangsfrequenz auf die Hälfte, was durch ein entsprechend niedrigeres Teilverhältnis oder Erhöhung der Eingangsfrequenz ausgeglichen werden kann. Auch komplette Frequenzteiler werden als integrierte Schaltkreise angeboten.

Abschließend sei die Nutzung des Zählerschaltkreises D 192 als Frequenzteiler mit wählbarem Teilverhältnis erläutert. Der Übertragsimpuls wird dabei zur Voreinstellung des Zählers und damit zu einer Einschränkung des Zählumfan-

ges genutzt. Auf Bild 5.87 wird der Zähler rückwärtszählend betrieben. Der Übertragsimpuls rückwärts lädt im Beispiel eine 7 in den Zähler ein. Da der Übertragsimpuls rückwärts normalerweise vom 1-0-Übergang der 0 bis zum 0-1-Übergang der 9 anliegt, durch das Laden die 9 aber übersprungen wird, verschwindet der

Übertragsimpuls sofort nach erfolgter Ladung der 7. Der folgende 0-1-Übergang an C2 setzt den Zähler damit auf die 6. Das Tastverhältnis am Ausgang Q3 beträgt $t_i:T = 1:2$, die Frequenzteilung $f_1:f_2 = 7:1$. Andere Teilverhältnisse sind durch Laden einer anderen Zahl möglich.



c)

Bild 5.86. Synchronfrequenzteiler 5:1 mit JK-Flip-Flop

a) Schaltung; b) Impulsdigramm; c) Tafel der logischen Werte

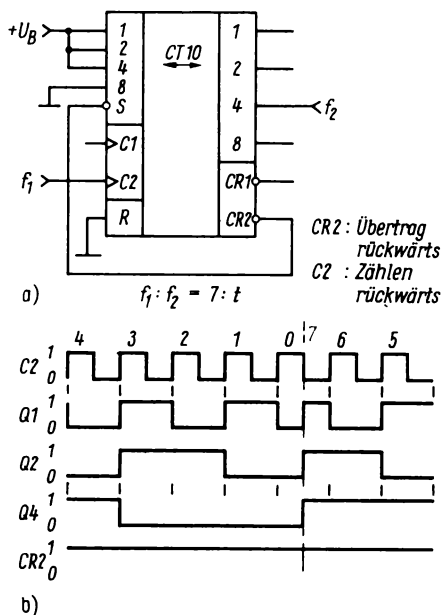


Bild 5.87. Synchronfrequenzteiler 7:1 mittels Dezimalzählerschaltkreises

a) Schaltung; b) Impulsdigramm

5.5.6.

Schieberegister und Schieberring

Ein Schieberegister mit vier Ausgängen (Bild 5.88) schiebt eine eingegebene Information mit jedem Taktimpuls um eine Stelle weiter. Fügt man Rückführungsleitungen vom Ausgang zum Eingang ein (Bild 5.88 rot eingetragen), so entsteht ein Schieberring. Eine Vorein-

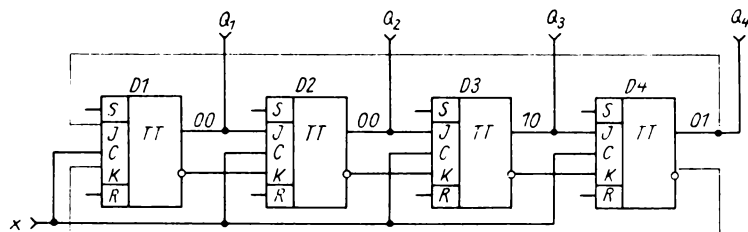


Bild 5.88. Schieberegister und Schieberring

stellung ist über die bei den einzelnen Flip-Flop vorhandenen R- und S-Eingänge möglich. Zur Erklärung der Vorgänge sei angenommen, daß durch die Voreinstellung an $Q1$, $Q2$ und $Q4$ Nullzustand und an $Q3$ Einszustand vorliegen (im Bild 5.88 an erster Stelle an den Ausgängen eingetragen). Da die Q-Ausgänge jeweils mit den J-Eingängen der folgenden Flip-Flop verbunden sind, kann beim folgenden Zählimpuls nur der Flip-Flop an Q auf Einszustand schalten, dessen J-Eingang vorher Einszustand hatte. Im Beispiel trifft das für den Flip-Flop D4 zu; Q4 kann auf Einszustand schalten. Alle anderen Stufen hatten am K-Eingang Einszustand, weshalb ihre Ausgänge Nullzustand behalten oder annehmen.

Das Schieberegister ist als Zähler verwendbar mit einem Zählumfang gleich der Anzahl der Einzelstufen. Eine Dekodierung ist nicht notwendig; dafür steigt der Aufwand an Flip-Flop gegenüber den oben beschriebenen Zähl-schaltungen. Mehrstufige Schieberegister werden ebenfalls als IS angeboten. Die Anzahl der Flip-Flop wird in der Bezeichnung angegeben (z. B. D 191C: 8-bit-Schieberegister; U 311D: 5-bit-Schieberegister).

Die integrierten Schieberegister unterscheiden sich sowohl durch die Anzahl der Ausgänge (bit) als auch durch die mögliche Schieberichtung und die Möglichkeiten der Eingabe. Diese kann parallel oder seriell erfolgen. Bei der Paralleleingabe ist für jede Stufe ein Voreinstelleingang erforderlich. Mit einem Ladebefehl wird der gewünschte Zustand innerhalb eines Taktes an allen Ausgängen gleichzeitig eingeladen. Weniger Anschlüsse werden bei serieller Eingabe benötigt. Hier wird der gewünschte Zustand von links in einer der Stufenzahl entsprechenden Taktzahl eingeschoben. Nachteilig ist die gegenüber der Paralleleingabe größere Voreinstellzeit.

Bei einem Schaltkreis mit 14 Anschlüssen kann man ein 4-bit-Schieberegister mit Paralleleingabe realisieren (4 Ausgänge, 4 Dateneingänge, 4 Befehlseingänge, 2 Stromversorgung), hingegen ist ein 8-bit-Schieberegister nur noch mit serieller Eingabe realisierbar.

5.5.7.

BCD-zu-7-Segment-Dekoder

An den Ausgängen der im Abschn. 5.5.4. vorgestellten Zählerschaltungen liegt der Zählerstand dual vor. Vielfach ist es aber zweckmäßig, den

Zählerstand dezimal anzuzeigen. Gut geeignet sind hierzu 7-Segment-Anzeigen. Die Aufgabe des Dekoders besteht nun darin, über eine kombinatorische Verknüpfung die eingegebene binär codierte Dezimalzahl (BCD) so zu dekodieren und neu zu kodieren, daß die jeweils zur Anzeige erforderlichen Segmente angesteuert werden. Im Bild 5.89 ist sowohl die Segmentkennzeichnung als auch die Zifferndarstellung der Ziffern 0 bis 9 für den TTL-Schaltkreis D 146/D 147 dargestellt. Am Beispiel des Segments „e“ soll ein Teil der inneren Schaltung erläutert werden (Bild 5.90a). Aus der Darstellung Bild 5.88 folgt, daß das Segment „e“ bei den Ziffern 0, 2, 6 und 8 anzeigen muß und bei den Ziffern 1, 3, 5, 7, 9 und 4 nicht anzeigen darf. Dazu werden die vier Eingangssignale mit den Wertigkeiten 1, 2, 4 und 8 (gekennzeichnet

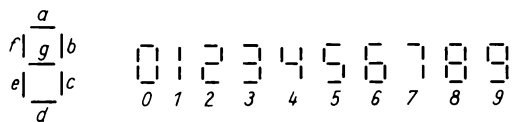


Bild 5.89. Siebensegmentdarstellung des Dekoders D 146/D 147

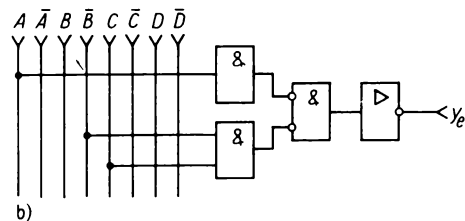
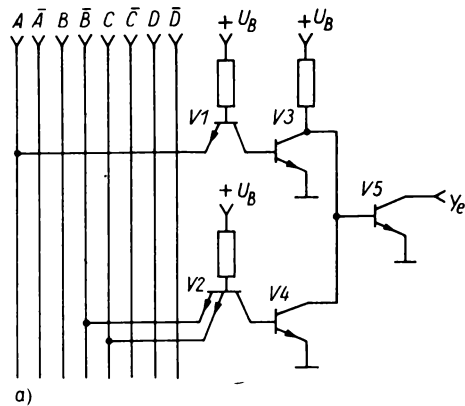


Bild 5.90. Teilschaltung des BCD-zu-7-Segment-Dekoders für das Segment „e“

a) innere Schaltung; b) logische Struktur der inneren Schaltung

an den Sammelschienen mit A, B, C und D) zunächst zusätzlich negiert (Sammelschienen $\bar{A}, \bar{B}, \bar{C}, \bar{D}$). Da der Dekoder über einen Ausgang mit offenem Kollektor verfügt, kann die Ansteuerung des Segments nur durch Stromaufnahme erfolgen, also bei durchsteuertem Transistor $V5$ (entspricht Pegel L am Ausgang). $V5$ ist nur dann durchsteuert, wenn die zueinander parallel geschalteten Transistoren $V3$ und $V4$ gesperrt sind, was wiederum nur möglich ist, wenn an den Emittern der vorgeschalteten Transistoren $V1$ und $V2$ Pegel L liegt.

Durch Pegel H am Emitter von $V1$ wird $V3$ für alle ungeradzahligten Zahlen geöffnet, da dann am Anschluß A Pegel H liegt. Damit bei der Zahl 4 Segment „e“ nicht zur Anzeige kommt, sind die Emitter von $V2$ an B und C angeschlossen, an denen dann Pegel H liegt. Für alle anderen Zahlen sind $V3$ und $V4$ gesperrt, $V5$ folglich durchsteuert. Bild 5.90b zeigt die logische Struktur der Teilschaltung für das Segment „e“. Die als IS gefertigten Dekoderschaltkreise können sowohl über den oben beschriebenen open collector output als auch über *Konstantstromausgänge* verfügen – meist als *Konstantstromsenken* ausgebildet. Die LED kann dann ohne zusätzlichen Strombegrenzungswiderstand direkt zwischen positive Betriebsspannung und Schaltkreisausgang geschaltet werden. Teilweise ist die Höhe des Konstantstromes über einen zusätzlichen Eingang beeinflussbar, was z. B. eine bequeme Anpassung der Helligkeit der Anzeige an die Umgebungshelligkeit gestattet.

Weitere Eingänge gestatten zusätzliche Funktionen. Dazu gehören der *Lampentest* (Aktivierung aller Ausgänge zur Kontrolle der Funktionsfähigkeit der Anzeige), die *Dunkeltastung* und die Dunkeltastung von Vornullen.

Verschiedene IS unterscheiden sich auch durch die Dekodierung der Dualzahlen 10 bis 15, deren Eingangszustände für weitere Anzeigen nutzbar sein können (Überlaufanzeige, Vorzeichenanzeige, Buchstabenanzeige zur Darstellung von Hexadezimalzahlen mit geringstmöglichem Aufwand).

Neben den zu-7-Segment-Dekodern werden auch Dekoder gefertigt, bei denen jeweils ein Ausgang entsprechend der Wertigkeit der Eingänge aktiviert wird (*1 aus n-Dekoder*).

5.5.8.

Multiplexer

Häufig ist man zur Verringerung der Anzahl benötigter Verbindungsleitungen oder wegen der konstruktiv begrenzten Anzahl der „pin“ von Schaltkreisen gezwungen, an mehreren Signalleitungen liegende Signale über nur eine Signalleitung zeitlich nacheinander zu übertragen (*Parallel-Serien-Wandlung*). Bild 5.91 zeigt, wie die an vier Eingangsleitungen parallel, d. h. zeitlich gleichzeitig liegenden Signale über eine kombinatorische Verknüpfung unter Zuhilfenahme eines Schieberings zeitlich nacheinander (seriell) auf den Ausgang geschaltet werden.

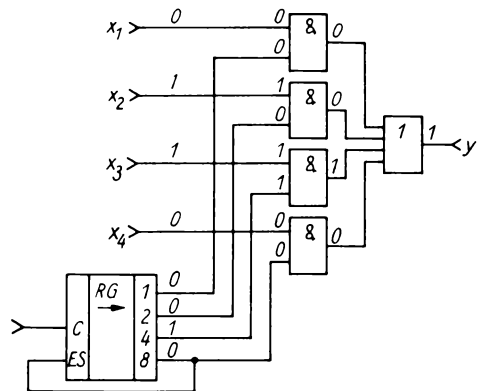


Bild 5.91. Prinzip eines Multiplexers mit Taktsteuerung

Es kann immer nur das Eingangssignal am Ausgang erscheinen, an dessen zugehörigem UND-Gatter am zweiten Eingang gerade der logische Wert 1 des Schieberings liegt. Dabei muß der Schiebering stets so voreingestellt sein, daß nur an einem Ausgang der logische Wert 1 liegt. Im Bild 5.91 sind beliebige logische Werte an x_1 bis x_4 eingetragen, wobei durch den Zustand des Schieberings gerade der an x_3 liegende logische Wert nach y übertragen wird.

Soll ein beliebiges Durchschalten der logischen Werte der Eingänge nach y möglich sein, sind die vier vom Schiebering kommenden Leitungen als Adreßleitungen zu nutzen und entsprechend der gewünschten Abfrage jeweils eine Leitung auf den logischen Wert 1 zu bringen.

Ein komplexes Beispiel soll die Zusammenhänge nochmals verdeutlichen. Dabei soll der Zählerstand eines hier nur dreistellig angenommenen Dezimalzählers durch eine dreistellige 7-Segment-Anzeige dargestellt werden. In der Schaltung nach Bild 5.92 sind für die Signalleit-

tungen der ersten Stelle durchgehende Linien, für die der zweiten Stelle Strichlinien und für die der dritten Stelle Strichpunktlinien verwendet worden. Die Stellentreiber der zweiten und dritten Stelle sind ebenso wie die Segmenttreiber der Segmente „a“ bis „f“ nicht dargestellt, da sie den gezeichneten Teilschaltungen des Stellentreibers der ersten Stelle und des Segmenttreibers für Segment „g“ entsprechen. Als Beispiel sind die logischen Werte für den Zählerstand 6 der Stelle 1 eingetragen. Die Verwendung des Multiplexers ermöglicht, daß für alle drei Stellen nur ein Dekoder erforderlich ist und nur sieben Segment- und drei Katodenleitungen zum Anzeigetableau zu führen sind. Gleiche Segmentanoden der einzelnen Ziffern sind miteinander verbunden und an den Dekoder über den Segmenttreiber angeschlossen. Die

einzelnen geführten Katodenleitungen ermöglichen, daß jeweils nur die Ziffer anzeigt, deren Katode über den Stellentreiber auf Masse geschaltet ist.

Ohne zeitmultiplexe Übertragung wären 3 Dekoder, 21 Segment- und eine Katodenleitung erforderlich. Der Schiebering wird mit hoher Frequenz getaktet (z. B. 1 kHz), weshalb die zeitlich nacheinander erfolgende Ansteuerung der drei Stellen der Ziffernanzeige vom Auge nicht wahrgenommen werden kann. Stellentreiber und Segmenttreiber sind darüber hinaus Beispiele für die Ansteuerung systemfremder Lasten.

Zusammenfassung

Sollen digitale Schaltungen über Kontakte angesteuert werden, sind Schaltmaßnahmen zur

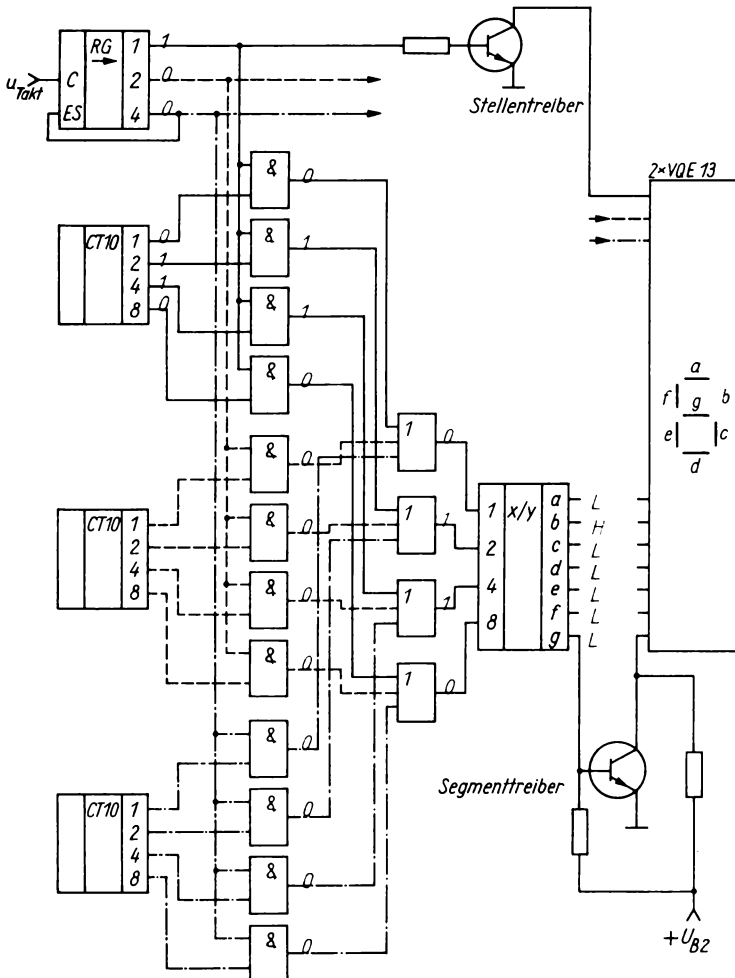


Bild 5.92. Zeitmultiplexe Ziffernansteuerung

Verhinderung der Auswirkung des Kontaktprelens erforderlich. Geeignet sind Flip-Flop, Integrierglieder mit nachgeschaltetem Schwellwertelement oder monostabile Elemente.

Notwendige Impulsverzögerungen erreicht man mit Kettenschaltung mehrerer Gatter oder mit Integriergliedern und nachgeschaltetem Schwellwertelement. Sind die Anstiegs- und Abfallzeiten eines Signals zu groß geworden, ist eine Flankenversteilerung durch Nachschaltung zweier Inverter oder durch ein Schwellwertelement möglich.

Zur Ansteuerung systemfremder Lasten sind Interface-Schaltungen erforderlich. Bei TTL-Schaltkreisen als Treiber sollte bevorzugt deren Stromaufnahmefähigkeit genutzt werden.

Dual- und Dezimalzähler können sowohl asynchron wie auch synchron arbeiten. Neben JK-Flip-Flop sind auch D-Flip-Flop und RS-Flip-Flop bei Zählern verwendbar. Eine modernere Lösung stellen die Zählerschaltkreise dar. Um einen Zählerstand mit einer 7-Segment-Anzeige darstellen zu können, gibt es entsprechende Dekoderschaltkreise. Die Dekodierung kann auch im Zählerschaltkreis mit realisiert sein.

Bei einem Schieberegister wird mit jedem Zählimpuls der jeweilige Ausgangszustand der einzelnen Flip-Flop um eine Stelle weiter geschoben. Anwenden kann man die Schieberegister für einfache Zähler und für Parallel-Serien-Wandler. Auch in der Datenverarbeitung finden sich vielfältige Anwendungsbeispiele.

Frequenzteiler nutzt man besonders zur Frequenzaufbereitung. Es lassen sich beliebige Teilverhältnisse erreichen. Die Ausgangsspannung hat im Regelfall ein von 1:2 abweichendes Tastverhältnis. Ein Tastverhältnis $t:T = 1:2$ ist erreichbar durch Nachschalten eines Zähl-Flip-Flop. Frequenzteiler werden als IS angeboten; sie sind teilweise programmierbar.

Zur Einsparung von Leitungen oder erforderlichen Anschlüssen an Schaltkreisen dienen Multiplexer, die ein an mehreren Leitungen parallel anliegendes Signal seriell über eine Leitung übertragen.

5.6.

Analog-Digital-Wandler und Digital-Analog-Wandler

5.6.1.

Allgemeines

Die Digitaltechnik hat durch die rasche Entwicklung der Mikroelektronik sowohl in der Informationsverarbeitung (z. B. durch Computer) als auch in der Informationsübertragung (z. B. Pulsmodulation, PCM) und in der Meßtechnik hohe Bedeutung erlangt. Bei genauer Betrachtung der genannten Beispiele erkennt man, daß die Eingangsgrößen und, mit Ausnahme der Meßtechnik, auch die Ausgangsgrößen jeweils analoger Natur sind. Daraus folgt die Notwendigkeit der Analog-Digital-Wandlung.

Innerhalb eines Prozesses kann z. B. die Drehzahl als zu erfassende analoge Größe auftreten. Damit der Computer die jeweilige Drehzahl verarbeiten kann, muß sie über einen Sensor erfaßt und einem ADU (*Analog-Digital-Umsetzer*) zugeführt werden. Das digitale Ergebnis der Informationsverarbeitung durch den Rechner muß nun im Allgemeinen durch einen DAU (*Digital-Analog-Umsetzer*) in eine analoge Größe rückgewandelt werden, um dann über einen Treiber als Stellgröße wieder die Drehzahl beeinflussen zu können.

Ähnlich sind die Verhältnisse auch bei der digitalen Informationsübertragung. Hier tritt beispielsweise die Sprache als analoges Signal auf. Das Mikrophon wandelt in eine entsprechende analoge Spannung, die durch einen ADU zu digitalisieren ist. Das entstandene digitale Signal kann nun übertragen werden und wird am Empfangsort durch einen DAU wieder in ein analoges Signal zurückverwandelt. Der Vorteil dieser doppelten Umsetzung liegt in einer wesentlich höheren Störsicherheit der Informationsübertragung (siehe Abschnitt 5.1. und Kapitel 6).

Die digitale Messung analoger Größen erfordert selbstverständlich auch einen ADU. Bei der digitalen Messung ist als Vorteil besonders die höhere Anzeigegenauigkeit gegenüber analog anzeigenden Meßinstrumenten zu werten. Kann man bei einem Zeigerinstrument allenfalls eine Ablesegenauigkeit von 1 % erreichen (im Meßbereich 1 V wäre als kleinste Wertänderung etwa 10 mV ablesbar), so liegt die Ablesegenauigkeit bereits bei einer 3-stelligen digitalen Anzeige bei 0,1 % (bei einem Meßbereich

von 1 V wäre die kleinste ablesbare Wertänderung 1 mV).

Tafel 5.4 zeigt Beispiele für die erreichbare Auflösung in Abhängigkeit von der Stufenzahl sowohl von Binär- als auch von BCD-Wandlern. Zu beachten ist dabei, daß die Auflösung allein nicht als Maß für die Güte eines ADU gewertet werden darf, da mögliche Fehler größer sein können als die Auflösung.

Tafel 5.4. Auflösung von ADU

	Quantisierungsstufen	Auflösung, bezogen auf den Endwert
6 bit	$2^6 = 64$	1,6 %
10 bit	$2^{10} = 1024$	0,098 %
14 bit	$2^{14} = 16384$	0,0061 % = 61 ppm (parts per million)
18 bit	$2^{18} = 262144$	3,8 ppm
$2 \frac{1}{2}$ Digit	200	0,5 %
3 Digit	1000	0,1 %
4 Digit	10000	0,01 % = 100 ppm

Bild 5.93 zeigt die Quantisierung einer analogen Größe durch einen dreistufigen Binärwandler. Die Quantisierungsstufen werden durch das niederwertigste bit (LSB; engl., *least significant bit*) bestimmt und der Endwert der digitalisierbaren analogen Größe mit FS (engl., *full scale*) bezeichnet. Das höchstwertige bit (MSB; engl., *most significant bit*) hat eine Wertigkeit von $\frac{1}{2}$ FS. Der Quantisierungsfehler beträgt $\pm \frac{1}{2}$ LSB, im Beispiel entsprechend $\pm \frac{1}{16}$ FS. So wird auf Bild 5.93 der Digitalwert 001 in einem Bereich der analogen Eingangsgröße von

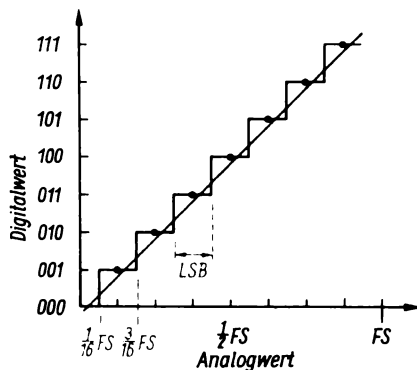


Bild 5.93. Quantisierung einer analogen Größe

$\frac{1}{16}$ FS bis $\frac{3}{16}$ FS ausgegeben oder bei Festlegung einer Spannung von 8 V für FS würde der genannte Digitalwert für einen Spannungsbereich von 0,5 V bis 1,5 V ausgegeben werden. Das Bild 5.93 zeigt auch, daß der praktisch erreichbare Endwert um 1 LSB niedriger als der nominelle Endwert ist.

Bei BCD-Wandlern bezeichnet man das niederwertigste Digit mit LSD (engl., *least significant digit*) und das höchstwertige mit MSD (engl., *most significant digit*).

Im Wandler können verschiedene Fehler auftreten. Bild 5.94 zeigt mit Kennlinie 2 den durch den Verstärker bedingten Offsetfehler (zur Erzielung der Nullanzeige ist ein bestimmter analoger Eingangswert erforderlich). Durch einen Nullabgleich kann dieser Fehler kompensiert werden. Kennlinie 3 zeigt auf dem gleichen Bild einen Verstärkungsfehler (auch Steilheits- oder Endwertfehler). Dieser Fehler wird durch einen Endwertabgleich kompensiert. Da beide Fehler praktisch immer gemeinsam vorliegen, wird sowohl ein Nullpunkt- als auch ein Endwertabgleich notwendig – es entsteht nach diesem Zweipunktabgleich die Wandlerkennlinie 4. Neben den genannten Fehlern können auch Linearitätsfehler auftreten, die z. B. im Ergebnis eines Zweipunktabgleiches (Wandlerkennlinie 4) entstehen oder durch Nichtlinearitäten der Wandlerschaltung. Je nach Größe dieses Fehlers verringert sich die tatsächlich nutzbare Auflösung des Wandlers.

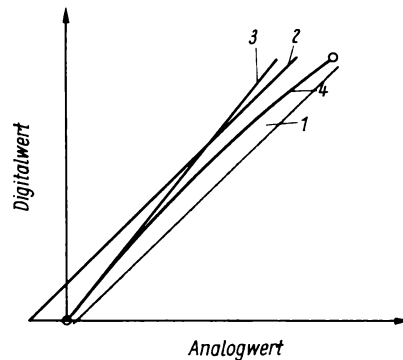


Bild 5.94. Wandlerkennlinie eines ADU

1 ideale Wandlerkennlinie; 2 Wandlerkennlinie mit Offset; 3 Wandlerkennlinie mit Verstärkungsfehler; 4 Wandlerkennlinie mit Zweipunktabgleich

Von Bedeutung für den jeweiligen Anwendungsfall ist auch die Umsetzzeit (engl., *conversion time*). Es ist verständlich, daß für die Um-

setzung eine bestimmte Zeit benötigt wird. Eine andere Beschreibung dieses Sachverhaltes stellt die *Umsetzrate* dar (Umsetzungen je Sekunde). Die realisierte Umsetzrate geht von 1 bis etwa 10^8 Messungen je Sekunde. Dabei läßt sich bei langsamen Wandlern gut eine Störspannungsunterdrückung durch mögliches integrierendes Verhalten erreichen. Derartige Wandler sind gut geeignet, wenn die Erfassung des Meßwertes durch den Menschen erfolgt (Ablesung). Sollen schnell ablaufende Vorgänge z. B. durch einen Rechner erfaßt werden, sind hingegen Wandler mit hoher Umsetzrate erforderlich.

Je nach Wandlerschaltung darf sich der Eingangswert während der Wandlung noch verändern oder nicht verändern. Im zweiten Fall ist es erforderlich, den Eingangswert erst abzutasten, zu speichern und danach zu wandeln (engl., sample and hold).

5.6.2.

Verfahren zur Analog-Digital-Wandlung

Bedingt durch die Vielfalt der Eigenschaften von ADU gibt es auch vielfältige Schaltungen. Tafel 5.5 zeigt eine mögliche Einteilung. Dabei haben die Zählverfahren die niedrigste und der Parallelumsetzer die höchste Umsetzrate. Der Schaltungsaufwand ist beim Parallelumsetzer am höchsten.

Verhältnismäßig einfach ist der Aufbau eines *Sägezahnumsetzers*. Sollen nur unipolare Spannungen gemessen werden, ist eine vereinfachte dargestellte Schaltung nach Bild 5.95 verwendbar. Kernstück dieser Schaltung ist ein Sägezahn-generator mit hoher Linearität des Anstieges. Dieser Sägezahn-generator wird von einer Steuerlogik gestartet und seine Ausgangsspan-

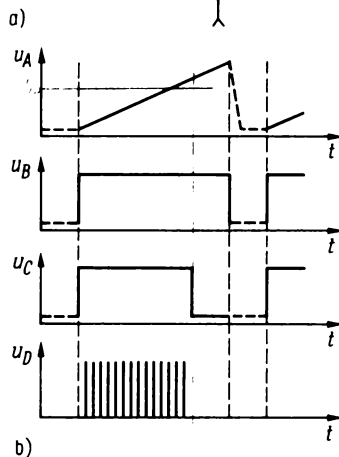
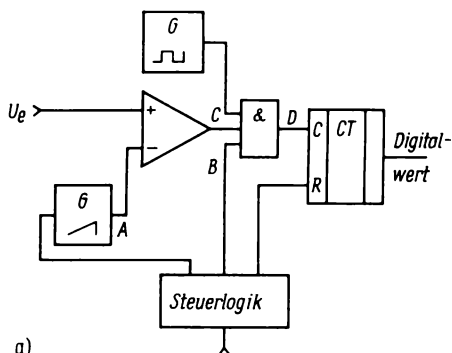
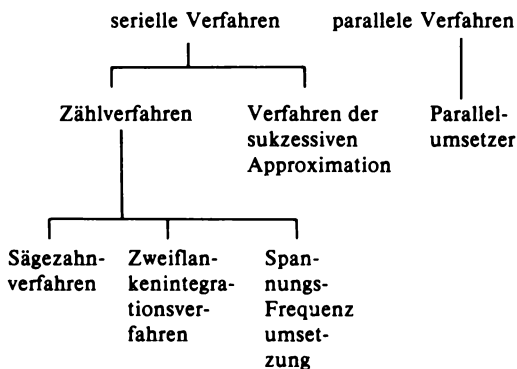


Bild 5.95. Unipolare Umsetzung nach dem Sägezahnverfahren

a) Schaltung; b) Ablaufdiagramm

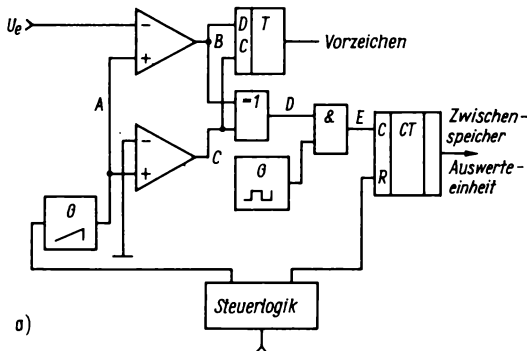
Tafel 5.5. Verfahren der Analog-Digital-Wandlung



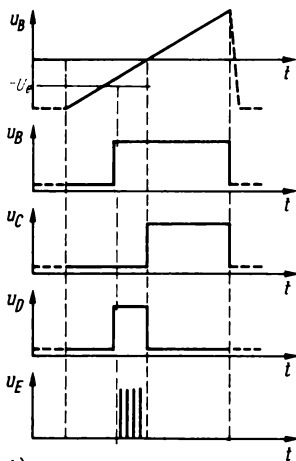
nung über einen Komparator mit der zu wandelnden Spannung verglichen. Solange die Sägezahnspannung noch kleiner als die Eingangsspannung ist, wird über das UND-Element die Rechteckspannung eines quartzstabilen Generators auf den Eingang des Zählers geschaltet. Erreicht die Sägezahnspannung nun den Wert der Eingangsspannung, so schaltet der Ausgang des Komparators um. Die Rechteckspannung gelangt nicht mehr zum Zähler. Der Zählerstand ist proportional der Zeit, die die Sägezahnspannung bis zum Erreichen der Eingangsspannung benötigte (*Spannungs-Zeit-Umsetzung*). Soll ein erneuter Meßvorgang ausgelöst werden, wird der vorherige Meßwert in einen hier nicht dargestellten Speicher übernommen, der Zähler durch die Steuerlogik auf 0 gesetzt und ein erneuter Sägezahn gestartet. Während der Meßpause wird das UND-Element durch die Steuerlogik gesperrt und zum Beginn des Sägezahnanstieges wieder freigegeben. Zur Auslösung des

Meßvorganges dient der Befehlseingang der Steuerlogik.

Das beschriebene Verfahren eignet sich bei geringfügiger Erweiterung auch zur Umsetzung bipolarer Spannungen. Hierzu ist dann ein Sägezahngenerator erforderlich, dessen Spannung von einem negativen Wert bis zu einem positiven Wert zeitlinear ansteigt. Der Zählvorgang läuft bei negativer Eingangsspannung vom Zeitpunkt der Gleichheit zwischen Sägezahn- und Eingangsspannung bis zum Nulldurchgang der Sägezahnspannung (Bild 5.96b) und bei positiver Eingangsspannung vom Nulldurchgang der Sägezahnspannung an bis zur Gleichheit zwischen Eingangs- und Sägezahnspannung. Die von der Polarität der Eingangsspannung abhängige Aufeinanderfolge des Schaltens der beiden Komparatoren (Meßkomparator und Nullpunktkomparator) wird zur Bildung des Vorzeichens genutzt. Im vorliegenden Fall schaltet der Meß-



a)



b)

Bild 5.96. Bipolare Umsetzung nach dem Sägezahnverfahren

a) Schaltung; b) Ablaufdiagramm

komparator vor dem Nullpunktkomparator – der Flip-Flop wird auf negatives Vorzeichen gesetzt.

Die Eingangsspannung sollte während der Messung möglichst konstant bleiben. Störspannungen beeinflussen die Umsetzung. Die Umsetzungsgenauigkeit hängt von der Langzeitkonstanz der den Sägezahnanstieg bestimmenden Bauelemente und von der Frequenzkonstanz des Zäufrequenzgenerators ab. Für hohe Auflösungen ist die Schaltung ungeeignet.

Nur eine Kurzzeitkonstanz von R und C ist hingegen beim Zweiflankenintegrationsverfahren erforderlich (engl., *dual slope integrating AD-converter*). Außerdem werden durch das integrierende Verhalten Störspannungen gut unterdrückt.

Beim Zweiflankenintegrationsverfahren wird zunächst ein Kondensator durch die zu wandelnde Spannung U_e innerhalb einer feststehenden Zeit zeitlinear aufgeladen. Im Bild 5.97 ist die Aufladekurve während der Aufladezeit t_M für zwei verschiedene Eingangsspannungen dargestellt.

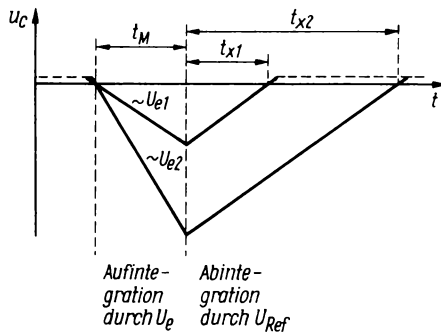


Bild 5.97. Ladespannungsverlauf beim Zweiflankenintegrationsverfahren

Mit dem Nulldurchgang der Ladespannung beginnt der Zählvorgang und der Ablauf der Zeit t_M . Die Aufladung wird beim Erreichen eines vorher festgelegten Zählerendstandes beendet (man kann mit einem bestimmten Zählerstand beginnen und durch den Übertrag die Ladung enden lassen oder mit 0 beginnen und nach Erreichung eines bestimmten Zählerstandes enden lassen). Jetzt schließt sich eine zeitlineare Entladung durch eine Referenzspannung an. Die Zeit bis zum Nulldurchgang der Spannung wird ausgezählt (im Bild 5.97 die Zeiten t_{x1} und t_{x2} für die zwei unterschiedlichen Eingangsspannungen) und als Digitalwert ausgegeben. Bild

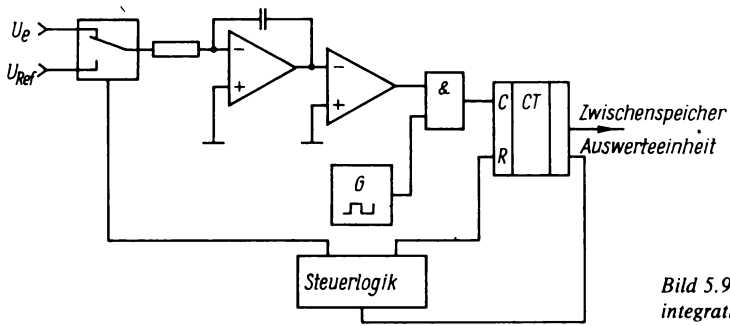


Bild 5.98. Prinzip des Zweiflankenintegrationsverfahrens

5.98 zeigt das vereinfachte Schaltungsprinzip. Für die zeitlineare Auf- und Entladung ist hier ein Miller-Integrator verwendet. (Denkbar ist an dieser Stelle auch die Wandlung der U_e in einen proportionalen Konstantstrom zur Aufladung des Kondensators und die Entladung durch eine Konstantstromquelle.) Der nachfolgende Komparator öffnet beim Nulldurchgang der Spannung das Tor, die Zählimpulse des Generators gelangen zum Zähler. Über die Steuerlogik wird die Abfolge der oben beschriebenen Vorgänge gesteuert.

Der Vorteil dieser Schaltung besteht darin, daß nur während des Wandlungsintervalls die die Auf- und Entladung bestimmenden Elemente konstante Werte haben müssen. So würde eine Erniedrigung der Zeitkonstante sowohl die Ladenspannung am Ende der Aufladezeit erhöhen als auch die Entladung steiler verlaufen lassen. Beide Faktoren kompensieren sich. Auch der Zählfrequenzgenerator braucht nur während des Wandlungsintervalls Frequenzkonstanz aufzuweisen. Bei einer Frequenzerhöhung würde zwar t_M und damit t_x verkürzt, was aber wegen der höheren Frequenz wieder eine gleiche Impulszahl während t_x ergeben würde. Diese Kurzzeitstabilität ist mit geringem Aufwand erreichbar. Lediglich die Konstanz der Referenzspannung muß auch über längere Zeit garantiert sein.

Die Wandlungszeit ist verfahrensbedingt groß – für schnelle Wandler ist das Verfahren ungeeignet. Gut geeignet hingegen ist es für Meßwertanzeigensysteme. Der Schaltkreis C 520 (in I²L-Technologie aufgebaut) arbeitet nach diesem Verfahren.

Eine Erweiterung des Zweiflankenintegrationsverfahrens stellt das *Mehrfankenintegrationsverfahren* dar. Hier wird ein weiterer Wandlungszyklus zur automatischen Nullpunkt Korrektur eingefügt (z. B. Wandlerfamilie C 500).

Für mittelschnelle bis schnelle Wandler eignet sich das Verfahren der *sukzessiven Approximation* (Bild 5.99). Hierbei wird über einen Komparator schrittweise die Eingangsspannung mit der Spannung eines internen DAU verglichen. Der Ablauf der Umsetzung ist für ein Beispiel auf Bild 5.100 dargestellt. Zunächst wird U_e mit dem *MSB* ($\triangleq \frac{1}{2} U_{max}$) verglichen. Im vorliegenden Fall ist U_e größer – es wird *MSB* gesetzt. Jetzt wird zu *MSB* $\frac{1}{2}$ *MSB* addiert und wiederum verglichen. Da U_e kleiner als $1,5$ *MSB* ist, wird $\frac{1}{2}$ *MSB* nicht gesetzt und im nächsten

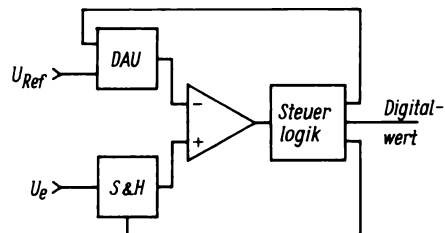


Bild 5.99. Prinzip des Verfahrens der sukzessiven Approximation

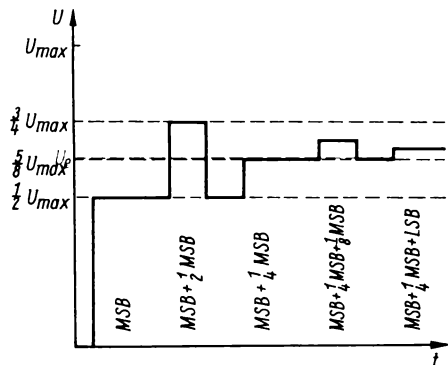


Bild 5.100. Ablaufplan beim Verfahren der sukzessiven Approximation

Schritt zu dem gesetzten $MSB \frac{1}{4} MSB$ addiert und verglichen. Jetzt ist U_e wieder größer, $\frac{1}{4} MSB$ wird gesetzt. Im folgenden Takt wird $\frac{1}{8} MSB$ addiert. Da dieser Wert wieder größer als U_e ist, wird $\frac{1}{8} MSB$ nicht gesetzt. Im letzten Takt wird LSB addiert und damit der Wert von U_e mit einer Abweichung kleiner $\pm \frac{1}{2} LSB$ erreicht. Die Wandlung ist beendet. Das Ergebnis lautet im Beispiel 10101.

Das Verfahren ist gegenüber Störspannungen größer als $\pm \frac{1}{2} LSB$ während der Umsetzzeit sowie gegen Änderungen der Eingangsspannung empfindlich. Es ist zweckmäßig, die zu wandelnde Spannung vor der Wandlung zu messen und zu speichern (sample & hold). Die benötigte Umsetzzeit ist meßwertunabhängig und richtet sich nur nach der Auflösung.

Die *Spannungs-Frequenz-Umsetzung* ist ebenfalls innerhalb von ADU verwendbar. Durch die U_e wird ein VCO (siehe Abschnitt 4.2.2.) gesteuert. Zählt man innerhalb einer feststehenden Meßzeit die Anzahl der Perioden des VCO, erhält man eine der Eingangsspannung proportionale Zahl.

Eine besonders hohe Umsetzrate ermöglicht der *Parallelwandler* (engl., flash converter). Für jede der vorgesehenen Quantisierungsstufen (bei einer Auflösung von 8 bit sind das 255) muß ein Komparator eingesetzt werden (Bild 5.101). An den einen Eingang aller Komparatoren wird die Eingangsspannung angelegt, an den anderen die entsprechend der Quantisierungsstufe geteilte

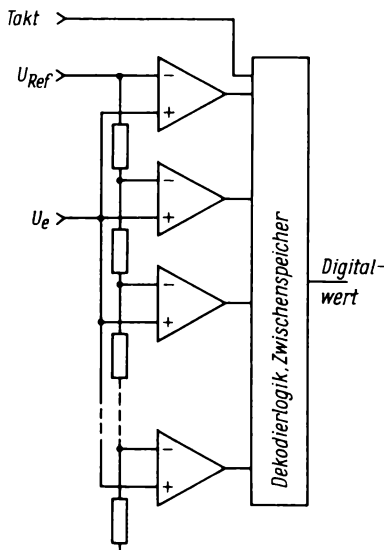


Bild 5.101. Prinzip des Parallelwandlers

Referenzspannung. Im Wandlungszeitpunkt wird festgestellt, bis zu welcher höchsten Quantisierungsstufe die Komparatoren geschaltet haben. Nach einer Dekodierung liegt der Digitalwert vor. Da für die Wandlung nur ein Takt benötigt wird, arbeitet dieser Wandler sehr schnell. Nachteilig ist der große Aufwand an Komparatoren. Hohe Auflösungen sind mit diesem Wandlerprinzip nicht zu erreichen.

(Einen Parallelwandler mit einer Quantisierung in 12 Stufen enthält der Schaltkreis A 277 – Ansteuerschaltkreis für LED-Zeilen).

5.6.3.

Verfahren zur Digital-Analog-Wandlung

Dem DAU kann der Digitalwert seriell oder parallel eingegeben werden. Bei serieller Eingabe werden die einzelnen Datenbit in Abhängigkeit von einem Takt über eine Eingangsleitung zeitlich nacheinander angelegt (ähnliche Verhältnisse liegen bei der PCM vor). Häufiger erfolgt eine *parallele Eingabe*, bei der alle Datenbit über jeweils eine eigene Leitung gleichzeitig dem DAU zugeführt werden. Der folgende Abschnitt bezieht sich nur auf DAU mit paralleler Eingabe.

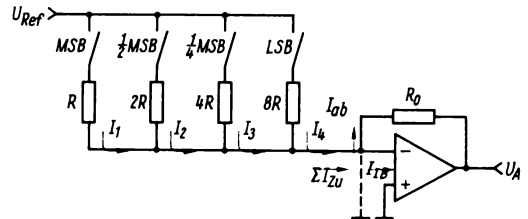


Bild 5.102. DAU mit gestuften Widerständen

Einen DAU mit gestuften Widerständen zeigt Bild 5.102. Zur Beschreibung der Wirkungsweise sollen zuerst die Verhältnisse am Operationsverstärker betrachtet werden. Wegen der sehr hohen inneren Verstärkung eines OV liegt in der vorliegenden Schaltung der invertierende Eingang virtuell auf Masse. Der in den invertierenden Eingang hineinfließende Strom ist wegen des sehr hohen Eingangswiderstandes äußerst gering und damit sehr viel kleiner als die anderen Ströme im Knotenpunkt S. Im Knotenpunkt S gilt deshalb

$$\Sigma I_{zu} = I_{ab}.$$

Für die Ausgangsspannung folgt weiter

$$U_A = -I_{ab} R_0 = -\Sigma I_{zu} R_0.$$

Die Ströme I_1 bis I_4 sind binär gestuft. So gilt z. B. $I_4 = 1/8 I_1$. Die Summe der binär gestuften Ströme im Knotenpunkt S hängt davon ab, welche Schalter beim jeweils zu wandelnden Digitalwert geschlossen sind.

Nachteilig bei dieser Schaltung ist, daß die Widerstände mit äußerst geringer Toleranz, geringem TK und hoher Alterungsbeständigkeit realisiert werden müssen. Bei höherer Auflösung kommt hinzu, daß der Wertebereich der Widerstände sehr groß wird. Bei 10 bit Auflösung müßte der zum *LSB* gehörige Widerstand 1024mal so groß sein wie der zum *MSB* gehörige Widerstand. Dieser Nachteil kann durch eine Gruppenbildung der Widerstände und zusätzliche Wichtung der Gruppen umgangen werden. Die konkrete Realisierung wird hier jedoch nicht betrachtet.

Eine große Bedeutung haben die $2R/R$ -Netzwerke in vielfältigen Schaltungsvarianten gefunden. Besonderer Vorteil dabei ist, daß nur zwei unterschiedliche Widerstandswerte benötigt werden. Bild 5.103 zeigt eine Schaltung mit einem solchen Netzwerk und nachfolgendem OV als Spannungsfolger. Entsprechend des zu wandelnden Digitalwertes wird die Stellung der Schalter festgelegt. In praktischen Schaltungen finden hier, wie auch in den vorangegangenen Schaltungen, natürlich elektronische Schalter Verwendung. Für den Fall, daß im Digitalwert nur das *MSB* vorliegt, ist die entstehende Widerstandsschaltung auf Bild 5.103b herausgezeichnet. Man erkennt, daß nach Reduzierung des Netzwerkes auf einen einfachen Spannungsteiler gilt

$$U_A = 1/2 U_{Ref}.$$

Für den Fall, daß sowohl *MSB* als auch $1/2$ *MSB* im Digitalwert enthalten sind (1100), ergibt sich ein Ersatzschaltbild gemäß Bild 5.103c. Zur Berechnung wird zunächst die untere Gruppe des Netzwerkes zusammengefaßt, U_x berechnet ($U_x = 5/8 U_{Ref}$), danach wird der Teiler der oberen Gruppe berechnet. Für U_A folgt

$$U_A = 3/4 U_{Ref}.$$

Abschließend soll ein DAW mit $2R/R$ -Netzwerk und nachgeschaltetem OV als Strom-Spannungs-Wandler betrachtet werden (Bild 5.104). Für den Knotenpunkt S gelten die bei der Schaltung mit gestuften Widerständen (Bild 5.102) gegebenen Erläuterungen. Die analoge Ausgangsspannung errechnet sich aus

$$U_A = -I_{ab} R_o = -\Sigma I_{zu} R_o.$$

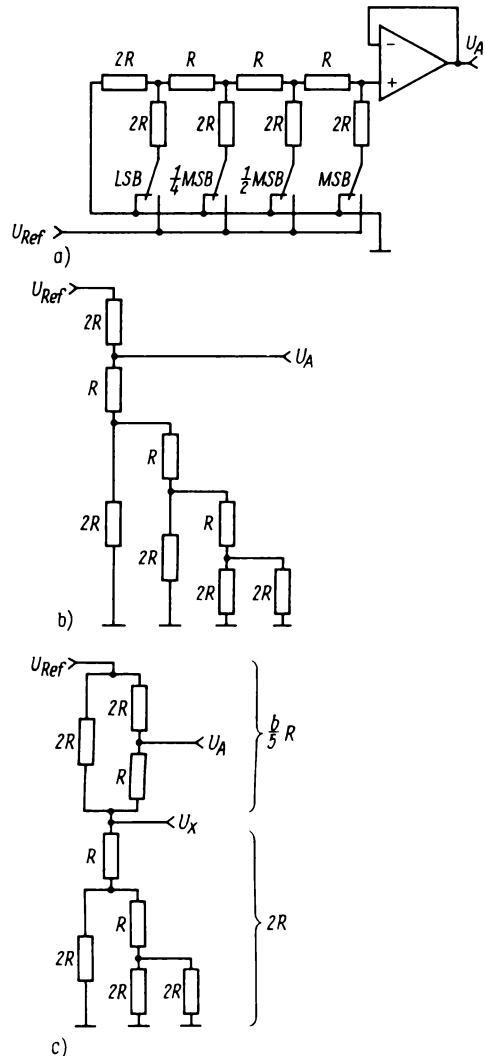


Bild 5.103. 4 bit DAW mit $2R/R$ -Netzwerk und Spannungsfolger

a) Schaltung; b) Widerstandsnetzwerk bei anliegendem *MSB*;
c) Widerstandsnetzwerk bei anliegendem *MSB* und $1/2$ *MSB* (1100)

Für den Fall, daß sowohl *MSB* als auch $1/2$ *MSB* gesetzt sind, ergibt sich die Teilschaltung nach Bild 5.104b. Im Knotenpunkt fließen die Ströme I_1 und I_2 zu. Für I_1 und I_2 folgt

$$I_1 = 1/2 U_{Ref}/R$$

$$I_2 = 1/4 U_{Ref}/R.$$

Damit gilt

$$\Sigma I_{zu} = I_{ab} = 3/4 U_{Ref}/R.$$

DAU werden auch als IS gefertigt. Vom Halbleiterwerk Frankfurt (Oder) wird dazu die DAU-Familie C 565 angeboten, die auf einer Erweiterung des Grundprinzips der gestuften und in gestufte Gruppen eingeteilten Widerstände basiert.

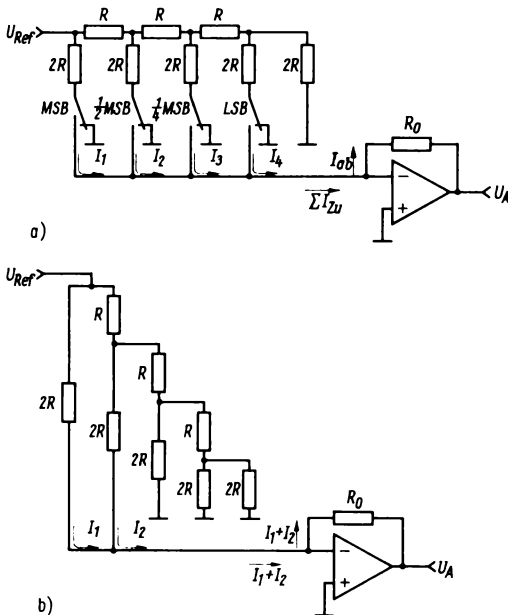


Bild 5.104. DAU mit $2R/R$ -Netzwerk und nachgeschaltetem Strom-Spannungs-Wandler

a) Schaltung; b) Ersatzschaltung für den Digitalwert 1100

Zusammenfassung

In der Informationsübertragung, Informationsverarbeitung und Meßtechnik wird heute vielfach digitale Schaltungstechnik eingesetzt. Ein- und Ausgangsgrößen sind meist analoger Natur. Deshalb sind sowohl ADU als auch DAU erforderlich. Die Wandler unterscheiden sich durch Auflösung und auftretende Fehler, Umsetzrate, Störspannungsempfindlichkeit, Anforderungen an die verwendeten Bauelemente, Möglichkeit der bipolaren oder unipolaren Umsetzung und den erforderlichen Aufwand. Bei den ADU werden serielle oder parallele Verfahren angewandt. Bei den DAU werden hauptsächlich Schaltungen mit gestuften Widerständen oder $2R/R$ -Netzwerken eingesetzt.

5.7.

Aufgaben

1. Warum ist die genaue Kenntnis der Innenschaltung digitaler Schaltungen für den Anwender nicht unbedingt erforderlich?
2. Geben Sie den TTL-Standard-Signalpegel an!
3. Informieren Sie sich über die in der DDR gefertigten Schaltkreisarten!
4. Welche Zuordnung besteht zwischen den Bezeichnungen H, L, 0, 1?
5. Erklären Sie die Wirkungsweise des Transistors als elektronischer Schalter!
6. Unter welcher Bedingung kann eine über der maximalen Verlustleistung liegende Leistung geschaltet werden (Herstellerangaben zum Transistor beachten)?
7. Welche besonderen Schaltungsmaßnahmen sind beim Schalten von Glühlampen und Induktivitäten zu beachten?
8. Welche kennzeichnenden Eigenschaften haben die vorstehend beschriebenen Grundgatter?
9. Erläutern Sie die Bedeutung des Lastfaktors bei TTL-Schaltkreisen!
10. Warum dürfen Ausgänge von TTL-Schaltkreisen mit „totem pole output“ nicht miteinander verbunden werden?
11. Was versteht man unter einer Transferkennlinie?
12. Erläutern Sie Eigenschaften und Anwendungsgebiete der drei möglichen Ausgänge von TTL-Schaltkreisen!
13. Welche Applikationshinweise folgen aus dem steilen Verlauf der Transferkennlinie (aktiver Bereich) von TTL-Schaltkreisen?
14. Warum ist bei TTL-Schaltkreisen eine stabilisierte Stromversorgung erforderlich?
15. Wodurch unterscheiden sich TTL-Standard und TTL-LS?
16. Wodurch unterscheiden sich Übersteuerungstechnik und Stromsteuertechnik?
17. Warum sollten unbenutzte Eingänge nicht unbeschaltet bleiben?
18. Was ist bei der Betriebsspannungszuführung zu beachten?
19. Beschreiben Sie die Wirkungsweise eines CMOS-, eines ECL- und eines I^2L -Gatters!
20. Warum sind CMOS-Schaltkreise für Anlagen der Landesverteidigung besonders geeignet?
21. Welche Aufgabe hat eine Interface-Schaltung? Wie kann sie aufgebaut werden?
22. Erläutern Sie die Möglichkeiten der Zusammenschaltung von CMOS, TTL-Standard und TTL-LS bzw., warum die direkte Zusammenschaltung nicht möglich ist!
23. Wodurch unterscheiden sich sequentielle und kombinatorische Schaltungen?
24. Beschreiben Sie die Eigenschaften der erläuterten Flip-Flop mittels Impulsdiagramms und Tafel der Übergänge!

25. Erläutern Sie an der logischen Struktur die Wirkungsweise eines RS- und eines statischen D-Flip-Flop!
26. Wie ist ein JK-Flip-Flop zu beschalten, um einen T-Flip-Flop zu erhalten?
27. Erläutern Sie die zeitliche Reihenfolge der Vorgänge bei einem Master-Slave-Flip-Flop in Abhängigkeit vom Taktimpuls!
28. Wie erreicht man eine hohe Sicherheit gegen Störimpulse?
29. Welche Eigenschaften hat ein Signalgenerator?
30. Erläutern Sie die Vorgänge innerhalb der Schaltung eines Signalgenerators!
31. Warum besteht in der Schaltung nach Aufgabe 30 die Gefahr einer Überschreitung der zulässigen Grenzwerte für die Eingangsspannung der Gatter?
32. Wie erreicht man ein Tastverhältnis von genau $t_1 : T = 1 : 2$?
33. Suchen Sie in der Fachliteratur weitere Signalgeneratorschaltungen!
34. Welche Eigenschaften hat ein monostabiles Element (Univibrator)?
35. Erläutern Sie die Vorgänge innerhalb der Schaltung eines monostabilen Gliedes! Welcher Zusammenhang besteht zwischen Erhol- und Verweilzeit?
36. Wodurch wird die Zeitdauer des Ausgangsimpulses bestimmt?
37. Bestimmen Sie die Eigenschaften weiterer in der Fachliteratur angegebener Schaltungen monostabiler Elemente! Achten Sie besonders darauf, ob in diesen Schaltungen die Grenzwerte für die Eingangsspannungen der Gatter nicht überschritten werden!
38. Welche weiteren Eigenschaften können als IS gefertigte monostabile Elemente aufweisen?
39. Erläutern Sie die Eigenschaften eines Schwellwertslements!
40. Wozu verwendet man ein Schwellwertslement?
41. Warum kann der Ausgang eines Schwellwertslements nur zwei diskrete Zustände annehmen?
42. Erläutern Sie andere Varianten der Schaltung von Schwellwertslementen!
43. Erläutern Sie Schaltungen, die die Auswirkungen des Kontaktprellens verhindern!
44. Unter welchen Bedingungen kann eine Impulsverzögerung erwünscht sein?
45. Wie erreicht man eine Impulsverzögerung?
46. Erläutern Sie Schaltungen zur Flankenversteilerung!
47. Wie können systemfremde Lasten an ein TTL-Gatter angeschlossen werden?
48. Entwerfen Sie Schaltungen, mit deren Hilfe eine Signalkleinlampe 12 V/50 mA durch einen TTL-Schaltkreis sowohl durch Stromaufnahme als auch durch Stromabgabe geschaltet werden kann! (Beachten Sie Abschn. 5.2.!)
49. Erläutern Sie die Funktion eines pull-up-Widerstandes! Wodurch wird sein Wert nach unten begrenzt?
50. Wodurch unterscheiden sich synchrone und asynchrone Zähler?
51. Erläutern Sie in der Schaltung nach Bild 5.79 die Vorgänge beim Übergang vom dritten zum vierten und vom sechsten zum siebenten Zählschritt!
52. Verfahren Sie für die Schaltung nach Bild 5.83 wie bei Aufgabe 51!
53. Erläutern Sie die Wirkungsweise des Synchronzählers!
54. Erläutern Sie Art und Funktion weiterer möglicher Ein- und Ausgänge von Dezimalzählerschaltkreisen!
55. Welche Aufgabe hat ein Schieberegister?
56. Erklären Sie die einzelnen Vorgänge beim Durchlauf eines 1-Signals durch den Schiebering!
57. Warum ist das Tastverhältnis des Teilers $1 : 10$ nicht $t_1 : T = 1 : 2$ in der Schaltung nach Bild 5.81?
58. Versuchen Sie, Frequenzteilerschaltungen mit anderen Teilverhältnissen zu entwerfen oder in der Fachliteratur vorgegebene Schaltungen zu interpretieren!
59. Wie erreicht man bei einem Frequenzteiler ein Tastverhältnis $t_1 : T = 1 : 2$?
60. Erläutern Sie die Beschaltung des D 192 als Frequenzteiler $5 : 1$ und $3 : 1$! Welches Tastverhältnis entsteht dabei? Betrachten Sie hierzu alle Ausgänge!
61. Welche Aufgabe hat ein Multiplexer?
62. Welche Aufgabe hat ein BCD-zu-7-Segment-Dekoder?
63. Welche Möglichkeiten der Ausgangsschaltung bei Dekodern gibt es? Erläutern Sie Eigenschaften, Vorteile und Anwendungsgebiete der jeweiligen Schaltung!
64. Erläutern Sie an Beispielen die Notwendigkeit der ADU und DAU!
65. Berechnen Sie die erreichbare Auflösung für einen 12 bit ADU und für einen $3\frac{1}{2}$ Digit ADU!
66. Vergleichen Sie die beschriebenen ADU hinsichtlich der erreichbaren Umsetzrate!
67. Erläutern Sie das Prinzip der unipolaren Umsetzung nach dem Sägezahnverfahren!
68. Warum ist beim Zweiflankenintegrationsverfahren nur eine Kurzzeitstabilität der Generatorfrequenz und von R und C notwendig?
69. Begründen Sie, warum das Verfahren der sukzessiven Approximation höhere Umsetzraten als die Zählverfahren ermöglicht!
70. Warum sind Parallelumsetzer nur für niedrige Auflösungen praktisch realisierbar?
71. Erläutern Sie das Prinzip des DAU mit gestuften Widerständen!
72. Begründen Sie an Hand der Eigenschaften des OV, warum bei der Schaltung gemäß Bild 5.102 gilt $U_A = -I_{ab} R_O$.

6.

Modulation und Demodulation

6.1.

Theoretische Grundlagen

6.1.1.

Allgemeines

In der Nachrichtentechnik versteht man unter Modulation die Beeinflussung einer Schwingung mit im allgemeinen relativ hoher Frequenz ω_h durch eine zeitabhängige Signal- oder Sendefunktion $s = f(t)$ (Sprache, Musik usw.), die im Vergleich zur Frequenz ω_h niedrige Frequenzen ω_n enthält. Die Sendefunktion wird damit aus ihrer ursprünglichen Frequenzlage in eine höhere umgesetzt. Diese Frequenzumsetzung ist für eine technisch realisierbare drahtlose Nachrichtenübertragung notwendig, weil nur für hohe Frequenzen Antennen mit geeignetem Strahlungswiderstand entworfen werden können und – das trifft auch für die drahtgebundene Nachrichtenübertragung zu – außerdem eine rationelle Mehrfachausnutzung der Übertragungskanäle ermöglicht wird.

Es kommt folglich darauf an, die Sendefunktion in einem Modulator einer hochfrequenten Trägerschwingung aufzuprägen, diese modulierte Trägerschwingung auszusenden sowie in einem geeigneten Medium zum Empfangsort zu leiten und dort zu demodulieren, d. h., die ursprüngliche Signal- oder Sendefunktion möglichst formgetreu wiederzugeben.

Die allgemeine Formulierung einer hochfrequenten Trägerschwingung lautet

$$u_h = \hat{U}_h \cos(\omega_h t + \varphi_h) \quad (6.1)$$

$\hat{U}_h$	Amplitude
φ_h	Winkel

Wegen der einfacheren mathematischen Umformungen wird eine cos- statt eine sin-Funktion zugrunde gelegt. Ihre Bestimmungsgrößen Amplitude $\hat{U}_h$ und Winkel $(\omega_h t + \varphi_h)$ können durch die Sendefunktion s beeinflusst werden, und man erhält damit folgende Modulationsarten:

● Amplitudenmodulation (AM)

Bei der AM wird die Amplitude $\hat{U}_h$ der Träger-

schwingung im Rhythmus der Sendefunktion s geändert, d. h. $\hat{U}_h = f(s)$.

● Winkelmodulation (WM)

Da der Winkel der Trägerschwingung sowohl die Frequenz ω_h als auch die Phase φ_h enthält, unterscheidet man bei der WM zwischen Frequenz- und Phasenmodulation.

Frequenzmodulation (FM)

Bei der FM wird die Frequenz ω_h der Trägerschwingung im Rhythmus der Sendefunktion s geändert, d. h. $\omega_h = f(s)$.

Phasenmodulation (PM)

Bei der PM wird die Phase φ_h der Trägerschwingung im Rhythmus der Sendefunktion s geändert, d. h. $\varphi_h = f(s)$.

Neben diesen Modulationsarten gibt es noch die Pulsmodulation. Bei dieser werden Impulsfolgen als Signal verwendet, z. B. Rechteck-, Dreieck-, Trapez- oder Kosinusimpulsfolgen. Eine Impulsfolge, bestehend z. B. aus Rechteckimpulsen, wird durch die Parameter Amplitude, Impulsdauer, Impulsfrequenz und Impulslage (Phase) bestimmt. Entsprechend dem beeinflussten Parameter ergeben sich folgende Arten:

● Pulsmodulation

Pulsamplitudenmodulation (PAM)

Pulsdauermodulation (PDM, auch Pulslängenmodulation PLM)

Pulsfrequenzmodulation (PFM)

Pulsphasenmodulation (PPM)

Pulskodemodulation (PCM)

Differenz-Pulskodemodulation (DPCM).

Die Bezeichnung der Modulationsarten PCM und DPCM ist nicht von den obengenannten Parametern abgeleitet worden.

Wie im Abschn. 3.6.1. angedeutet, läßt sich nach *Fourier* jede Funktion in eine Summe von Sinus- bzw. Kosinusfunktionen zerlegen. Folglich ist auch jede Sendefunktion aus Sinus- bzw. Kosinusschwingungen aufgebaut. Es ist daher möglich und im Interesse der Übersichtlichkeit sinnvoll, wenn im folgenden als einfachster Fall die Aussteuerung durch eine Sendefunktion mit einer einzigen Kosinusfunktion der Frequenz ω_n behandelt wird.

6.1.2.

Darstellung der Modulationsarten

Amplitudenmodulation

Bei der Amplitudenmodulation wird die Amplitude der hochfrequenten Trägerschwingung im Rhythmus der Sendeschwingung s verändert. Die Sendeschwingung sei dabei eine cos-förmige Spannung mit einer im Verhältnis zur Trägerschwingung niedrigen Frequenz ω_n der Form

$$s = \hat{U}_n \cos \omega_n t \quad (6.2)$$

Unter Annahme, daß die Phase der HF-Schwingung $\varphi_h = 0$, ergibt sich für die modulierte Schwingung u_{AM} die Zeitfunktion

$$u_{AM} = [\hat{U}_h + \hat{U}_n \cos \omega_n t] \cos \omega_h t \quad (6.3)$$

bzw. nach Ausklammern der Amplitude $\hat{U}_h$

$$u_{AM} = \hat{U}_h \left[1 + \frac{\hat{U}_n}{\hat{U}_h} \cos \omega_n t \right] \cos \omega_h t \quad (6.4)$$

Die maximale Amplitudenänderung der modulierten Trägerschwingung heißt Amplitudenhub $\hat{U}_n$, während die Intensität der Modulation durch den Quotienten $\hat{U}_n / \hat{U}_h$ bestimmt ist und Modulationsgrad m genannt wird:

$$m = \frac{\hat{U}_n}{\hat{U}_h} \quad (6.5)$$

Durch Auflösen der Klammer in Gl. (6.4) erhält man für die modulierte Spannung

$$u_{AM} = \hat{U}_h \cos \omega_h t + \hat{U}_n \cos \omega_h t \cos \omega_n t.$$

Unter Zuhilfenahme des Additionstheorems $2 \cos \alpha \cos \beta = \cos(\alpha + \beta) + \cos(\alpha - \beta)$ folgt für den letzten Summanden obestehender Gleichung

$$\hat{U}_n \cos \omega_h t \cos \omega_n t = \frac{\hat{U}_n}{2} \cos(\omega_h + \omega_n) t + \frac{\hat{U}_n}{2} \cos(\omega_h - \omega_n) t,$$

und man erhält als Zeitfunktion der modulierten Spannung

$$u_{AM} = \hat{U}_h \cos \omega_h t + \frac{\hat{U}_n}{2} \cos(\omega_h + \omega_n) t + \frac{\hat{U}_n}{2} \cos(\omega_h - \omega_n) t \quad (6.6a)$$

bzw. nach Einführung des Modulationsgrads m gemäß Gl. (6.5)

$$u_{AM} = \hat{U}_h \left[\cos \omega_h t + \frac{m}{2} \cos(\omega_h + \omega_n) t + \frac{m}{2} \cos(\omega_h - \omega_n) t \right] \quad (6.6b)$$

Die Gl. (6.6a) und (6.6b) liefern folgendes Resultat:

Die modulierte Schwingung u_{AM} enthält neben der Trägerschwingung mit der Frequenz ω_h und der Amplitude $\hat{U}_h$ zwei weitere Schwingungen mit den Frequenzen $\omega_h + \omega_n$ und $\omega_h - \omega_n$, die jeweils mit der halben Signalamplitude $\hat{U}_n/2$ auftreten. Letztere liegen beiderseits der Trägerfrequenz und heißen Seitenfrequenzschwingungen. Die Summenfrequenz $\omega_h + \omega_n$ wird als obere, die Differenzfrequenz $\omega_h - \omega_n$ als untere Seitenfrequenz bezeichnet. Beide enthalten die zu übertragende Nachricht.

Neben der analytischen sind die im Bild 6.1 angegebenen Darstellungen einer amplitudenmodulierten Schwingung geeignet. Bild 6.1c zeigt als Zeigerdiagramm die Verhältnisse in einem bestimmten Zeitpunkt. Auf dem sich mit der Frequenz ω_h drehenden Zeiger $\hat{U}_h$ der Trägerschwingung bewegen sich zwei Zeiger der Länge

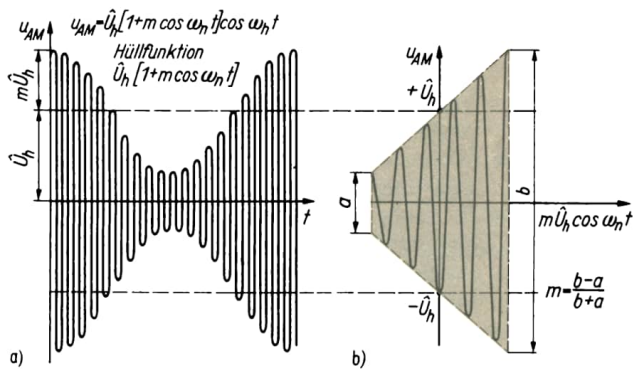
$$\frac{\hat{U}_n}{2} = \frac{m}{2} \hat{U}_h$$

mit der Frequenz ω_n , aber zueinander mit gegenläufiger Drehrichtung. Durch vektorielle Addition der drei Zeiger erhält man die Amplitude der modulierten Schwingung im betrachteten Zeitpunkt. Der resultierende Zeiger der beiden Seitenfrequenzzeiger ist stets in Phase bzw. Gegenphase mit dem Trägerfrequenzzeiger. Die Amplitude der modulierten Schwingung schwankt damit zeitlich zwischen den Werten $\hat{U}_h + m\hat{U}_h$ und $\hat{U}_h - m\hat{U}_h$.

Bild 6.1a zeigt den Zeitverlauf von u_{AM} im Liniendiagramm, Bild 6.1b als Modulationstrapez. Letzteres entsteht, wenn der Elektronenstrahl eines Oszillografen in x-Richtung durch die Signalspannung

$$u_n = \hat{U}_n \cos \omega_n t = m \hat{U}_h \cos \omega_n t$$

und in y-Richtung durch die modulierte Spannung u_{AM} abgelenkt wird. In dieser Weise wird häufig der Modulationsgrad gemessen; denn es gelten für die Seite $a = 2(\hat{U}_h - \hat{U}_n) = 2\hat{U}_h(1 - m)$ und für die Seite



resultierender Zeiger der beiden Seitenfrequenzzeiger

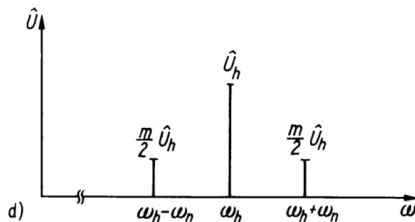
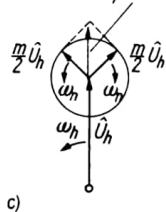


Bild 6.1. Darstellung einer kosinusförmig amplitudenmodulierten Hochfrequenzschwingung

a) Liniendiagramm; b) Modulationstrapez; c) Zeigerdiagramm; d) Spektraldiagramm

$b = 2 (\hat{U}_h + \hat{U}_n) = 2 \hat{U}_h (1 + m)$. Damit ist

$$\frac{b-a}{b+a} = \frac{2 \hat{U}_h (1+m) - 2 \hat{U}_h (1-m)}{2 \hat{U}_h (1+m) + 2 \hat{U}_h (1-m)} = \frac{4 \hat{U}_h m}{4 \hat{U}_h} = m.$$

Das Spektraldiagramm nach Bild 6.1d zeigt, daß neben der Trägerschwingung mit der Amplitude $\hat{U}_h$ und der Kreisfrequenz $\omega_h = 2\pi f_h$ zwei Seitenschwingungen mit der Amplitude $\frac{m}{2} \hat{U}_h$ und den Kreisfrequenzen $\omega_h + \omega_n = 2\pi(f_h + f_n)$ und $\omega_h - \omega_n = 2\pi(f_h - f_n)$ auftreten. Die Bandbreite b des Spektrums ergibt sich daraus mit

$$b = 2\omega_n. \quad (6.7)$$

Nun sollen die Leistungen berechnet werden, die bei einer amplitudenmodulierten Schwingung im Träger und in den Seitenfrequenzschwingungen enthalten sind. Man erhält diese, wenn man die in Gl. (6.6) bzw. Bild 6.1 dargestellten Amplituden quadriert. Es ergibt sich damit als Trägerleistung

$$P_h = \text{konst.} \cdot \hat{U}_h^2 \quad (6.8)$$

und Leistung einer Seitenfrequenzschwingung

$$P_{(h \pm n)} = \text{konst.} \cdot \hat{U}_h^2 \frac{m^2}{4} \quad (6.9)$$

Die Gesamtleistung einer AM-Schwingung ist

$$P_{AM} = P_h + 2P_{(h \pm n)} = P_h \left(1 + 2 \frac{P_{(h \pm n)}}{P_h} \right) = P_h \left(1 + \frac{m^2}{2} \right) \quad (6.10)$$

Wie Gl. (6.9) zeigt, kann die Leistung einer Seitenfrequenzschwingung selbst bei 100%iger Modulation nicht mehr als 25 % der Trägerleistung betragen. Die Gesamtleistung P_{AM} steigt dann entsprechend Gl. (6.10) auf den 1,5fachen Wert der Leistung der unmodulierten Schwingung (Bild 6.2).

In der Praxis ist es aber im allgemeinen notwendig, daß nicht nur mit einer Signalfrequenz, sondern mit einem ganzen Signalfrequenzband moduliert wird. Es ergibt sich dann ein in Regellage oberes und ein in Kehrlage unteres Seitenfrequenzband, und die Gesamtleistung der AM-Schwingung kann dabei maximal auf das Vierfache der Leistung der unmodulierten

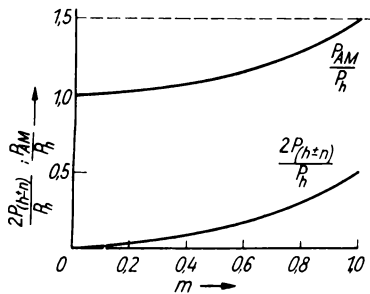


Bild 6.2. Leistung einer AM-Schwingung in Abhängigkeit vom Modulationsgrad

Schwingung anwachsen. Bei Rundfunksendern beträgt der Modulationsgrad im Mittel aber nur $m = 30\%$, weil bei höherer Modulationsintensität im Modulator Verzerrungen auftreten und außerdem die Amplitude der Signalfrequenzschwingung je nach Lautstärke schwankt. Damit kann abschließend festgestellt werden, daß die Amplitudenmodulation eine schlechte Leistungsbilanz aufweist. Die Bandbreite des Frequenzspektrums ist offensichtlich von der maximalen Signalfrequenz abhängig und beträgt nach Gl. (6.7)

$$b = 2\omega_{n\max}.$$

Im Interesse der rationellen Ausnutzung der Übertragungskanäle wurde der Trägerfrequenzabstand für Rundfunksender im Lang- und Mittelwellenbereich mit 9 kHz festgelegt, so daß die Signalfrequenz

$$f_n = \frac{\omega_n}{2\pi} = \frac{9 \text{ kHz}}{2} = 4,5 \text{ kHz}$$

am Kanalende liegt.

Die obigen Darlegungen gelten für die Zweiseitenband-Amplitudenmodulation mit Träger. Dabei steckt nur ein geringer Teil der von einem Sender abgestrahlten Energie in den Seitenbändern, die jedes für sich die gesamte Nachricht enthalten. Zur Nachrichtenübermittlung genügt es also, wenn nur ein Seitenband übertragen wird. Auch die Trägerschwingung, die nur zur Demodulation erforderlich ist, kann ganz oder teilweise unterdrückt und empfängerseitig wieder zugesetzt werden.

Die Einseitenband-Amplitudenmodulation EAM mit Trägerrest hat große praktische Bedeutung. Trotz höheren technischen Aufwands hat die EAM gegenüber der ZAM folgende Vorteile:

- Der Bandbreitenbedarf ist nur halb so groß.
- Die erforderliche Sendeleistung ist geringer, bzw. die vorhandene Leistung wird durch Verteilung auf nur ein Seitenband besser ausgenutzt.
- Wegen der geringeren Bandbreite ist das Empfängerrauschen kleiner.
- Es lassen sich Übermodulation und starke Verzerrungen, die als Folge des in der Ionosphäre auftretenden selektiven Trägerschwunds entstehen können, vermeiden.

Wegen dieser Vorteile verwendet man die EAM in der Trägerfrequenztechnik und im kommerziellen Funkbetrieb; EAM mit Träger wird in der Fernsehtechnik benutzt.

Winkelmodulation

Die Winkelmodulation läßt sich in Frequenz- und Phasenmodulation einteilen.

Frequenzmodulation

Bei der Frequenzmodulation wird die Frequenz der hochfrequenten Trägerschwingung im Rhythmus der Sendeschwingung s verändert. Diese Sendeschwingung sei wiederum cos-förmig mit einer im Verhältnis zur Trägerschwingung niedrigen Frequenz ω_n der Form

$$s = \Delta\Omega \cos \omega_n t \quad (6.11)$$

Unter Annahme, daß die Phase der HF-Schwingung $\varphi_h = 0$, ergibt sich für die modulierte Schwingung u_{FM} folgende Zeitfunktion:

$$u_{FM} = \hat{U}_h \cos [\omega_h + \Delta\Omega \cos \omega_n t] t \quad (6.12)$$

Die maximale Frequenzänderung der modulierten Trägerschwingung heißt Frequenzhub $\Delta\Omega$, während die Intensität der Modulation durch den auf die Sendefrequenz ω_n bezogenen Frequenzhub bestimmt ist und Modulationsindex x genannt wird:

$$x = \frac{\Delta\Omega}{\omega_n} \quad (6.13)$$

Der Frequenzhub $\Delta\Omega$ entspricht der Amplitude der Sendefunktion, ist also maßgebend für die Lautstärkeunterschiede. Er ist von der Sendefrequenz selbst unabhängig; denn diese bestimmt lediglich die Häufigkeit, mit der die Frequenz des Trägers vom kleinsten zum größten Wert innerhalb einer Sekunde schwankt (Bild 6.3).

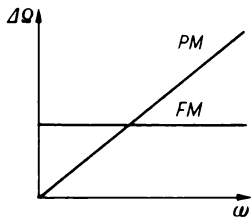


Bild 6.3. Frequenzhub der Frequenz- und Phasenmodulation in Abhängigkeit von der Sendefrequenz

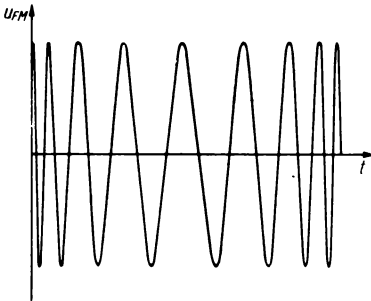


Bild 6.4. Liniendiagramm einer FM-Schwingung

Der Modulationsindex x hingegen ist gemäß Gl. (6.3) bei gleichbleibendem Frequenzhub umgekehrt proportional der Sendefrequenz, d. h., daß mit höher werdenden Signalfrequenzen die Intensität der Frequenzmodulation sinkt. Das Liniendiagramm im Bild 6.4 zeigt den Verlauf einer frequenzmodulierten Schwingung entsprechend Gl. (6.12).

Die oben angegebene Gl. (6.12) kann man weiterentwickeln, wenn man bedenkt, daß der in der Klammer stehende Ausdruck eine von ω_n abhängige Kreisfrequenz und das gesamte Argument $[\omega_h + \Delta\Omega \cos \omega_n t] t$ ein Phasenwinkel ist.

Über den allgemeinen Zusammenhang, daß die Kreisfrequenz oder Winkelgeschwindigkeit in einem bestimmten Augenblick der Differentialquotient des Winkels nach der Zeit ist, liefert eine Integration der Kreisfrequenz über die Zeit den Winkel. Diese Rechenoperation ist mit Hilfe der höheren Mathematik sehr einfach, soll aber hier nicht durchgeführt werden. Sie liefert aus Gl. (6.12) das Resultat

$$u_{FM} = \hat{U}_h \cos \left[\omega_h t + \frac{\Delta\Omega}{\omega_n} \sin \omega_n t \right]$$

(6.14 a)

bzw. mit Gl. (6.13)

$$u_{FM} = \hat{U}_h \cos [\omega_h t + x \sin \omega_n t] \quad (6.14 b)$$

Die durch Gl. (6.14) angegebene Zeitfunktion einer FM-Schwingung läßt sich durch das Zeigerdiagramm in Bild 6.5 leicht interpretieren.

Wie bei der AM-Schwingung bewegen sich auch hier zueinander gegenläufig auf der Spitze des sich mit der Kreisfrequenz ω_h drehenden Zeigers $\hat{U}_h$ der Trägerschwingung die den beiden Seitenfrequenzschwingungen zugeordneten Zeiger mit der Kreisfrequenz ω_n . Im Gegensatz zu Bild 6.1 c steht aber der aus den beiden Seitenfrequenzzeigern resultierende senkrecht auf dem Pfeil $\hat{U}_h$, da die Amplitude $\hat{U}_{FM}$ konstant bleiben muß. Das wird allerdings auch nur erreicht, wenn sich die Spitze des Zeigers $\hat{U}_{FM}$ auf dem Bogen $a-b$ und nicht auf der Geraden $A-B$ bewegt; letzteres würde eine gleichzeitig auftretende Amplitudenmodulation bedeuten. Eine konstante Amplitude kann aber nur durch weitere umlaufende Zeiger erklärt werden, die Seitenfrequenzschwingungen höherer Ordnung repräsentieren. Aus Bild 6.5 ist zu erkennen, daß um so mehr Seitenfrequenzen höherer Ordnung auftreten werden, je größer die Phasenänderung ist. Die Amplituden der Trägerfrequenz- und Seitenfrequenzschwingungen 1. und höherer Ordnung lassen sich mit Hilfe der von dem deutschen Mathematiker Bessel aufgestellten Bessel-Funktionen berechnen.

Bild 6.6 zeigt ein mit diesen Funktionen ermitteltes Spektraldiagramm für einen Modulationsindex $x = 5$.

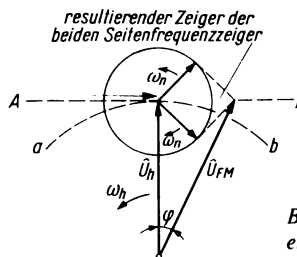


Bild 6.5. Zeigerdiagramm einer FM-Schwingung

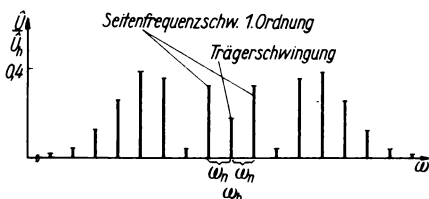


Bild 6.6. Spektraldiagramm einer FM-Schwingung

Die Spektrallinien der Seitenfrequenzschwingungen liegen beiderseits symmetrisch zur Trägerschwingung und treten jeweils im Abstand der Modulationsfrequenz ω_n auf. Theoretisch entstehen unendlich viele Seitenfrequenzschwingungen; in der Praxis vernachlässigt man aber die Seitenschwingungen höherer Ordnung, deren Amplituden weniger als 3 bis 5 % der unmodulierten Trägeramplitude betragen.

Bei einer frequenzmodulierten Schwingung ist bei gleicher Aussteuerung und damit konstantem Frequenzhub der das Spektraldiagramm bestimmende Modulationsindex von der Modulationsfrequenz ω_n abhängig. Das bedeutet, daß um so mehr Seitenfrequenzen zu berücksichtigen sind, je kleiner die Signalfrequenz ist, und umgekehrt. Bei Signalfrequenzänderung ändert sich bei FM außer dem Abstand der Spektrallinien auch die Form des Spektraldiagramms. Wird beispielsweise die Modulationsfrequenz verdoppelt, dann sinkt gemäß Gl. (6.13) der Modulationsindex auf den halben Wert, z. B. von $x_1 = 4$ auf $x_2 = 2$. Bei einem kleineren Modulationsindex sind entsprechend weniger Seitenfrequenzschwingungen zu berücksichtigen als bei einem größeren. Da sich aber der Abstand der Spektrallinien ebenfalls verdoppelt, ändert sich die Bandbreite des Spektrums nur unwesentlich und ist in erster Linie vom Frequenzhub abhängig. In Näherung ist die Bandbreite des Spektrums

$$b \approx 2 [\Delta\Omega + (1 \dots 2) \omega_n] \quad (6.15 a)$$

Im Bild 6.7 werden die Verhältnisse anschaulich gezeigt. Im Bild 6.7 b wurde eine gegenüber Bild 6.7 a doppelt so hohe Modulationsfrequenz

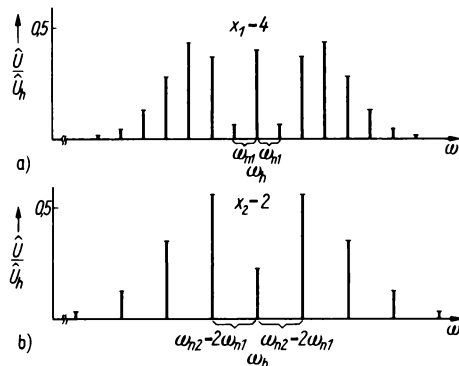


Bild 6.7. Spektraldiagramm einer FM-Schwingung mit einem Modulationsindex

a) $x_1 = 4$; b) $x_2 = 2$

gewählt, d. h. $\omega_{n2} = 2 \omega_{n1}$, die bei gleicher Aussteuerung einen halb so großen Modulationsindex $x_2 = \frac{1}{2} x_1$ zur Folge hat, z. B. $x_2 = 2$ gegenüber $x_1 = 4$.

Phasenmodulation

Bei der Phasenmodulation wird die Phase der hochfrequenten Trägerschwingung im Rhythmus der Sendeschwingung verändert. Diese Sendeschwingung sei wiederum cos-förmig mit einer im Verhältnis zur Trägerschwingung niedrigen Frequenz ω_n der Form

$$s = \Delta\Phi \cos \omega_n t \quad (6.16)$$

Unter Annahme, daß die Ausgangsphase der HF-Schwingung $\varphi_h = 0$, ergibt sich für die modulierte Schwingung u_{PM} folgende Zeitfunktion:

$$u_{PM} = \hat{U}_h \cos [\omega_h t + \Delta\Phi \cos \omega_n t] \quad (6.17)$$

Die maximale Phasenänderung der modulierten Trägerschwingung heißt Phasenhub $\Delta\Phi$, durch den die Intensität der Modulation bestimmt wird, die dem Modulationsindex der Frequenzmodulation analog ist. Der Phasenhub $\Delta\Phi$ entspricht der Amplitude der Sendefunktion, ist also maßgebend für die Lautstärkeunterschiede; die Sendefrequenz dagegen bestimmt die Häufigkeit, mit der der Phasenwinkel der Trägerschwingung in einer Sekunde hin- und herschwankt.

Stellt man die durch Gl. (6.17) angegebene Zeitfunktion der phasenmodulierten Schwingung im Liniendiagramm dar, dann erhält man die gleiche Form wie bei der Frequenzmodulation entsprechend Bild 6.4. Bei einer Modulation mit einer sin- bzw. cos-förmigen Schwingung lassen sich Phasen- und Frequenzmodulation nicht unterscheiden; erst bei Modulation mit zwei oder mehreren Frequenzen treten die Unterschiede hervor.

Im Vergleich der beiden Zeitfunktionen Gl. (6.14) und Gl. (6.17) erkennt man bis auf eine Verschiebung der Ausgangsphase von 90° (sin statt cos!) Übereinstimmung, wenn man

$$\frac{\Delta\Omega}{\omega_n} = \Delta\Phi \quad (6.18)$$

setzt.

Es ergibt sich für die Phasenmodulation folglich das gleiche Zeigerbild wie bei der Frequenzmodulation. Zu beachten ist dabei lediglich, daß

im Fall der FM für einen gleichbleibenden Frequenzhub der Phasenhub mit zunehmender Signalfrequenz kleiner wird und im Fall der PM für einen gleichbleibenden Phasenhub der Frequenzhub mit größer werdender Signalfrequenz ansteigt.

Daß sich Frequenz- und Phasenhub umgekehrt proportional verhalten, ist auch aus dem Zeigerdiagramm nach Bild 6.5 ersichtlich.

Unmoduliert erscheint lediglich der Zeiger $\hat{U}_h$. Wird nun mit einem bestimmten Frequenzhub $\Delta\Omega$ und zunächst niedriger Signalfrequenz ω_n ausgesteuert, dann wird die Augenblicksfrequenz relativ lange von der Ruhefrequenz ω_h der unmodulierten Schwingung abweichen. Da der Phasenhub aber der Phasenunterschied zwischen modulierter und unmodulierter Schwingung und damit durch die Differenz der Anzahl der Nulldurchgänge gegeben ist, wird bei höherer Signalfrequenz und gleichem Frequenzhub die Frequenzabweichung zwar ebenso groß sein, jedoch nur kürzere Zeit andauern. Damit hat der Zeiger $\hat{U}_{FM}$ am Ende des Frequenzhubs bereits nach einem kleineren Phasenwinkel die der Frequenz am Ende des Frequenzhubs entsprechende Geschwindigkeit.

Nach den obigen Darlegungen erkennt man, daß bei gleicher Lautstärke im Fall der Frequenzmodulation der Phasenhub mit wachsender Signalfrequenz ω_n abnimmt, dagegen bei Phasenmodulation der Frequenzhub zunimmt (Bild 6.3). Da bei der Phasenmodulation die hohen Signalfrequenzen, vom Standpunkt der Frequenzmodulation aus gesehen, bevorzugt werden, hat man beim wechselseitigen Überführen der Modulationsarten folgendes zu beachten: Durch ein vor den Modulationseingang geschaltetes Differenzierglied (Hochpaß) wird eine FM- in eine PM-Schwingung verwandelt; umgekehrt wird durch ein vorgeschaltetes Integrierglied (Tiefpaß) eine PM- in eine FM-Schwingung übergeführt. Eine gemeinsame mathematische Behandlung der Frequenz- und Phasenmodulation ist durch die wichtige Beziehung nach Gl. (6.18) gegeben, indem bei FM der von der Signalfrequenz ω_n abhängige Phasenhub $\Delta\Phi = \Delta\Omega/\omega_n$ bzw. bei PM der von der Signalfrequenz ω_n abhängige Frequenzhub $\Delta\Omega = \Delta\Phi\omega_n$ einzuführen ist.

Das Spektraldiagramm der PM-Schwingung stimmt bei gleichem Phasenhub mit dem der FM-Schwingung nach Bild 6.6 völlig überein. Im Gegensatz zur FM ist aber bei der PM die Amplitudenverteilung unabhängig von der

Höhe der Modulationsfrequenz ω_n . Wird ω_n verändert, dann bleibt bei der phasenmodulierten Schwingung wegen des konstant zu haltenen Phasenhubs die Form des Spektraldiagramms erhalten; es ändert sich lediglich der Abstand ω_n der einzelnen Spektrallinien, und das ganze Diagramm rückt mehr oder weniger zusammen. Die sich ergebende Bandbreite ist folglich nach der gleichen Beziehung Gl. (6.15a) zu ermitteln, nur daß anstelle $\Delta\Omega = \Delta\Phi\omega_n$ zu setzen ist.

$$b \approx 2 [\Delta\Phi\omega_n + (1 \dots 2) \omega_n] \quad (6.15b)$$

Die folgende Leistungsbetrachtung gilt sowohl für FM als auch für PM.

Da bei FM und PM die Amplitude der modulierten Schwingung konstant bleibt, wird auch die von einem Sender abgestrahlte Gesamtleistung nicht verändert:

$$P_{FM, PM} = P_h = \text{konst.} \quad (6.19)$$

Dabei ist P_h die Trägerleistung ohne Modulation. Bei Aussteuerung verändert sich diese und wird mehr oder weniger auf die Seitenfrequenzschwingungen verteilt. Es ergibt sich dann

$$P_{FM, PM} = P_h = P_0 + 2 P_{(h \pm n)}, \quad (6.20)$$

wobei P_0 die Trägerleistung mit Modulation angibt.

Diese ist von der Modulationsintensität $x = \Delta\Phi$ abhängig und durch die Bessel-Funktionswerte 0. Ordnung J_0 bestimmt:

$$P_0 = P_h J_0^2(x); \quad J_0 = g(\Delta\Phi). \quad (6.21)$$

Die Leistung der Seitenfrequenzschwingungen $2 P_{(h \pm n)}$ erhält man, wenn man von der Gesamtleistung die des modulierten Trägers, also P_0 , abzieht. Sie muß damit auch vom Phasenhub abhängig sein und ergibt

$$2 P_{(h \pm n)} = 2 P_h \sum_{p=1}^n J_p^2(x) \quad J_p = g(\Delta\Phi). \quad (6.22)$$

J_p sind die von $\Delta\Phi$ abhängigen Bessel-Funktionswerte 1., 2. bis n-ter Ordnung. Führt man die Gln. (6.21) und (6.22) in (6.20) ein, dann erhält man als Resultat für die Gesamtleistung

$$P_{FM, PM} = P_h \left[J_0^2(x) + 2 \sum_{p=1}^n J_p^2(x) \right] \quad (6.23)$$

Bild 6.8 zeigt die auf $P_{FM, PM}$ bezogene Seitenfrequenzleistung in Abhängigkeit vom Phasenhub $\Delta\Phi$. Man erkennt daraus, daß bei Phasenhüben $\Delta\Phi > 1,6$ die Leistung der Seitenfre-

quenzschwingungen stets über 80 % der Gesamtleistung liegt. Insbesondere bei den Werten für $\Delta\Phi \approx 2,4; 5,5; 8,7; 11,8; \dots$ [Nullstellen von $J_0(x)$] verschwindet der Träger, so daß die Gesamtleistung ausschließlich auf die Seitenfrequenzschwingungen verteilt ist. Die Leistungsbilanz von FM und PM ist also wesentlich besser als die der AM.

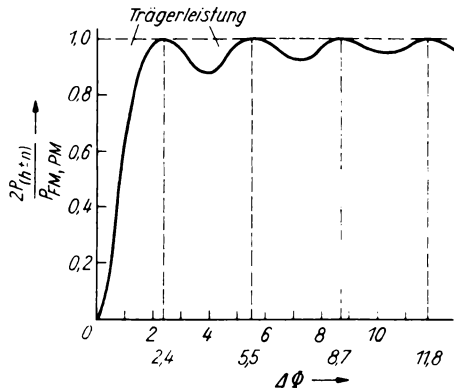


Bild 6.8. Leistung der FM/PM-Seitenschwingung in Abhängigkeit vom Phasenhub

Bei der Übertragung breitbandiger Signale wendet man in der Praxis weder die reine Frequenz- noch die reine Phasenmodulation an, sondern die *Frequenzmodulation mit Preemphasis*. Preemphasis ist eine lineare Vorverzerrung der Sendeschwingungen, wobei mittels eines Hochpasses die Schwingungen höherer Frequenzen angehoben werden. Damit geht die Frequenzmodulation bei höheren Frequenzen in eine Phasenmodulation über (Bild 6.9). Empfängerseitig muß die Preemphasis durch eine *Deemphasis* kompensiert werden. Das geschieht nach der Demodulation mittels eines Tiefpasses. Die Zeitkonstante des Hoch- und Tiefpasses beträgt beim Rundfunk im allgemeinen $50 \mu\text{s}$. (In anderen Regionen wird mit Zeitkonstanten gearbeitet, die geringfügig von den hier angegebenen abweichen.)

Durch Pre- und Deemphasis wird die Wiedergabequalität wesentlich verbessert. Das durch die

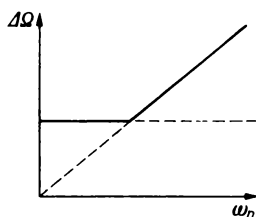


Bild 6.9. Frequenzhub der Frequenzmodulation mit Preemphasis in Abhängigkeit von der Sendefrequenz

Übertragung verursachte Rauschen steigt nämlich mit größer werdender Frequenz an. Wird nun durch die Deemphasis die angehobene Sendespannung wieder abgesenkt, dann wird damit auch proportional die Rauschspannung vermindert. Der Signal-Rausch-Abstand ist dann nahezu frequenzunabhängig und bei hohen Frequenzen nicht schlechter als bei tiefen.

Vermerkt sei noch, daß bei FM-Rundfunksendungen der maximale Frequenzhub

$$\frac{\Delta\Omega_{\max}}{2\pi} = 75 \text{ kHz}$$

bei einem Signalfrequenzband von

$$\frac{\omega_n}{2\pi} = 30 \text{ Hz} \dots 15 \text{ kHz}$$

beträgt.

Wegen der großen Bandbreite des Spektrums ist deshalb auch der Ausdruck „Breitbandfrequenzmodulation“ gebräuchlich.

Abschließend sollen ohne ausführliche Erläuterungen die wichtigsten Vor- und Nachteile der Frequenz- und Phasenmodulation gegenüber der Amplitudenmodulation zusammengestellt werden:

Vorteile:

- besserer Senderwirkungsgrad
- Verzerrungen durch gekrümmte Kennlinien der aktiven Bauelemente und Dämpfungsverzerrungen treten nicht auf
- die Modulation kann leistungsarm erfolgen
- wegen konstanter Amplitude $U_{\text{FM, PM}}$ keine Kreuzmodulation
- geringe Störungen durch Nachbarsender und aperiodische Störgeräusche, die im wesentlichen amplitudenmodulierter Art sind und im Empfänger durch Begrenzer unterdrückt werden.

Nachteile:

- Die erforderliche Bandbreite des Übertragungskanal muß größer als bei AM sein.

Pulsmodulation

Bei der Pulsmodulation wird eine *Signalformwandlung* vorgenommen. Hauptanwendungsbereich ist die Mehrfachausnutzung eines Übertragungskanal durch das Zeitmultiplexverfahren (mit Ausnahme der PFM).

Pulsamplitudenmodulation

Die zu wandelnde Sendeschwingung s wird nur zu bestimmten Zeiten abgetastet. Am Ausgang

des Abtastglieds entsteht eine Impulsfolge, bei der die Amplituden der Impulse dem Momentanwert der Sendeschwingung zu den Abtastzeiten entspricht (Bild 6.10). Je kleiner die Zeitabstände der Abtastimpulse (Taktimpulse) sind, desto genauer kann die Sendeschwingung nachgebildet werden. Für eine verzerrungsarme Modulation ist es notwendig, eine Periode der Sendeschwingung mindestens zweimal abzutasten, d. h., die Abtastfrequenz wird durch die höchste im zu wandelnden Signal enthaltene Frequenz bestimmt.

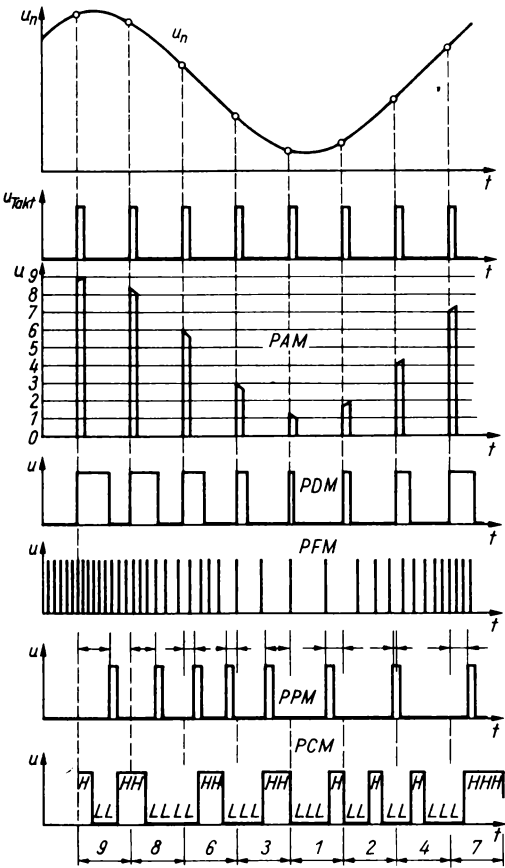


Bild 6.10. Pulsmodulationsarten

Pulsdauermodulation

Die zu wandelnde Sendeschwingung s wird in eine möglichst große Anzahl von Einzelimpulsen konstanter Amplitude, aber unterschiedlicher Dauer zerlegt (Bild 6.10). Die Impulsdauer jedes Einzelimpulses entspricht dem jeweiligen Momentanwert der Sendeschwingung.

Pulsfrequenzmodulation

Die zu wandelnde Sendeschwingung s wird in Impulse konstanter Amplitude und Dauer, aber unterschiedlicher Anzahl zerlegt (Bild 6.10). Die Impulsfrequenz entspricht dem jeweiligen Momentanwert der Sendeschwingung.

Pulsphasenmodulation

Die zu wandelnde Sendeschwingung s wird in Impulse konstanter Amplitude und Dauer, aber unterschiedlicher Lage (Phase) zerlegt (Bild 6.10). Die Impulsphase jedes Einzelimpulses entspricht dem jeweiligen Momentanwert der Sendeschwingung.

Pulscode-Modulation

Die zu wandelnde Sendeschwingung s wird bei der PCM zunächst *abgetastet* und in eine PAM-Impulsfolge übergeführt. Die Amplituden der Impulse werden dann *quantisiert*, die quantisierten Amplitudenwerte anschließend *kodiert*.

Im gewählten Beispiel (Bild 6.10) wurden die Impulsamplituden durch neun Quantisierungsstufen erfasst. Je mehr Quantisierungsstufen vorgesehen werden, desto genauer lassen sich die Impulsamplituden unterscheiden. Die wirklichen Abtastwerte werden also etwas fehlerbehaftet erfasst, und es entsteht ein unvermeidliches Quantisierungsrauschen. Um beispielsweise Sprache in ausreichender Qualität zu übertragen, benötigt man $2^7 = 128$ Quantisierungsstufen.

Die Impulsamplituden werden anschließend kodiert. Als Kode wird ein Binärkode mit den logischen Werten 1 und 0 verwendet. Für das angegebene Beispiel ist ein Viererkode erforderlich, d. h., jeder Wert lässt sich durch 4 bit darstellen.

2^3	2^2	2^1	2^0	
L	L	L	L	LLLL = 0
L	L	L	H	LL LH = 1
L	L	H	L	LL HL = 2
L	L	H	H	LL HH = 3
L	H	L	L	LH LL = 4
L	H	L	H	LH LH = 5
L	H	H	L	LH HL = 6
L	H	H	H	LH HH = 7
H	L	L	L	HL LL = 8
H	L	L	H	HL LH = 9

Bild 6.11. Bitfolgen für die Quantisierungswerte von 0 bis 9

Im Bild 6.11 sind die Bitfolgen für die Quantisierungswerte von 0 bis 9 zusammengestellt. Die Impulse haben die Pegel H und L. Im Bild 6.10 ist das PCM-Signal für die oben angegebene Sendeschwingung s dargestellt.

Differenz-Pulskodemodulation

Bei der DPCM wird lediglich die Abweichung zwischen dem tatsächlichen Abtastwert und einem Vorhersagewert übertragen. Man benötigt dadurch nur eine geringere Bandbreite als bei PCM. Beim Dekodieren müssen dann die Abweichung und der in einem Speicher befindliche Vorhersagewert addiert werden.

Mit Hilfe der Pulsmodulation ist es möglich, mehrere Sendefunktionen in einem gemeinsamen Übertragungskanal zu übertragen. Das dabei angewendete Verfahren heißt *Zeitmultiplexverfahren*. Bezogen auf die PAM (Bild 6.10) heißt das, daß in den Lücken zwischen den Impulsen weitere Impulse anderer Sendeschwingungen eingefügt werden. Man erhält dann ineinandergeschachtelte Impulsfolgen, deren zeitliche Zuordnung dem Empfänger durch zusätzliche Synchronisierungsimpulse mitgeteilt werden muß.

Die zeitmultiplexe Übertragung bei PCM erfolgt so, daß zunächst das erste Kodewort der ersten Sendeschwingung, dann das jeweils erste Kodewort der zweiten, dritten Sendeschwingung usw. erfaßt werden. Erst wenn das jeweils erste Kodewort aller zeitlich ineinandergeschachtelter Sendeschwingungen übertragen wurde, folgt das jeweils zweite Kodewort der ersten, zweiten, dritten Sendeschwingung usw. Die PFM ist für Zeitmultiplex nicht geeignet, weil sich die Impulsanzahl jeder Sendeschwingung ständig ändert. Die Impulse der verschiedenen Sendeschwingungen würden sich so ineinanderschachteln, daß sie empfängerseitig nicht mehr voneinander getrennt werden könnten. Die PFM wird aber häufig in der Fernwirk- und Steuerungstechnik verwendet.

Sollen pulsmodulierte Signale drahtlos übertragen werden, so ist eine Frequenzumsetzung durch Amplitudenmodulation oder Winkelmodulation erforderlich. Man unterscheidet zwischen der primär angewendeten Pulsmodulationsart und der *sekundär angewendeten* Modulationsart AM oder WM. Es ergeben sich daraus eine Reihe möglicher Kombinationen, von denen die Kombination PPM-AM besondere Bedeutung hat und in der Richtfunktechnik angewendet wird.

Bei der Frequenzumsetzung durch AM treten auch die HF-Impulse nur zu bestimmten diskreten Zeitpunkten und wegen der PPM nur mit konstanter Amplitude auf. Damit wird die Leistungsstufstufe des Senders günstig ausgenutzt. Da die Abtastimpulse insbesondere bei zeitmultiplexer Übertragung relativ schmal sein müssen, kann eine wirkungsvolle Frequenzumsetzung nur dann erfolgen, wenn die Periodendauer der HF klein gegenüber der Dauer der Abtastimpulse ist. Das führt in Verbindung mit der Forderung, die HF-Energie gerichtet abzustrahlen, zur Anwendung von Frequenzen ≥ 200 MHz.

Der große *Vorteil* der Puls- gegenüber der Amplituden- und Winkelmodulation liegt in der wesentlich störungsfreieren Nachrichtenübertragung. Besonders bei PCM muß nämlich empfängerseitig nur festgestellt werden, ob Impulse vorhanden sind oder nicht. *Nachteilig* ist allerdings der große Bandbreitenbedarf für die Übertragung.

Zusammenfassung

Eine der wesentlichen Aufgaben der Nachrichtentechnik ist es, ein Signal mit einer sin- oder cos-förmigen Zeitfunktion vor der Übertragung aus seiner ursprünglichen Frequenzlage in eine höhere umzusetzen. Dieser Vorgang ist durch das Wirken einer Trägerschwingung gekennzeichnet und heißt Modulation. Am Empfangsort hat eine mit Demodulation bezeichnete Rücktransformation zu erfolgen, um das ursprüngliche Signal möglichst formgetreu wiederzugeben.

Wird die Amplitude der Trägerschwingung im Rhythmus der Sendefunktion verändert, dann spricht man von *Amplitudenmodulation AM*. Die AM ist neben der Trägerschwingung durch das Auftreten von zwei die Nachricht enthaltenden Seitenfrequenzschwingungen gekennzeichnet, die die Bandbreite des Spektrums bestimmen, also $b = 2 \omega_n$.

Die Modulationsintensität wird durch das Amplitudenverhältnis von Signal- und Trägerschwingung gekennzeichnet und heißt Modulationsgrad ($m = \hat{U}_n / \hat{U}_b$).

Der Modulationsgrad wird meist in Prozent angegeben und muß stets ≤ 100 % sein.

Wird der Winkel der Trägerschwingung im Rhythmus der Sendefunktion verändert, dann spricht man von *Winkelmodulation WM*. Da die Winkeländerung sowohl durch eine Frequenz- als auch durch eine Phasenwinkeländerung er-

folgen kann, unterteilt man die WM in Frequenz- und Phasenmodulation.

Bei der *Frequenzmodulation FM* wird die Frequenz der Trägerschwingung im Rhythmus der Sendefunktion verändert. Der Maximalwert der Frequenzabweichung heißt Frequenzhub. Er entspricht der Amplitude der Signalfunktion und ist von deren Frequenz unabhängig.

Die FM hat theoretisch unendlich viele Seitenfrequenzschwingungen zur Folge. Unterschreiten ihre Amplituden 3 bis 5 % der Trägeramplitude, können sie unberücksichtigt bleiben. Man erhält dann eine Bandbreite des Spektrums von $b \approx 2[\Delta\Omega + (1 \dots 2) \omega_n]$.

Die Modulationsintensität wird durch den signalfrequenzabhängigen Modulationsindex $x = \Delta\Omega/\omega_n$ gekennzeichnet.

Bei der *Phasenmodulation PM* wird der Phasenwinkel der Trägerschwingung im Rhythmus der Sendefunktion verändert. Der Maximalwert der Phasenabweichung heißt Phasenhub. Er entspricht der Amplitude der Signalfunktion und ist von deren Frequenz unabhängig.

Phasen- und Frequenzmodulation können durch die Beziehung

$$\Delta\Phi = \frac{\Delta\Omega}{\omega_n} \Leftrightarrow \Delta\Omega = \Delta\Phi\omega_n$$

wechselseitig übergeführt werden.

Das Spektraldiagramm bei PM weist ebenfalls theoretisch unendlich viele Seitenfrequenzschwingungen auf. Im Gegensatz zur FM bleibt aber dabei die Amplitudenverteilung unabhängig von der Signalfrequenz. Es ändert sich je nach Größe der Modulationsfrequenz lediglich der Abstand der einzelnen Spektrallinien, und das ganze Spektrum rückt mehr oder weniger zusammen. Die Bandbreite ist also proportional der Signalfrequenz und über die Beziehung $\Delta\Omega = \Delta\Phi\omega_n$ mit der bei FM angegebenen Bandbreitengleichung zu ermitteln.

Die Modulationsintensität ist bei PM durch den Phasenhub bestimmt. In der Praxis wendet man meist die Frequenzmodulation mit Preemphase an.

Die Leistungsbilanz der Amplitudenmodulation ist schlecht. Selbst bei 100%iger Modulation beträgt die Leistung einer Seitenfrequenzschwingung nur 25 % der Trägerleistung. Die Leistungsbilanz der Winkelmodulation, also FM und PM, ist gut. Schon bei Phasenhuben $\Delta\Phi > 1,6$ liegt die Leistung der Seitenfrequenzschwingungen stets über 80 % der Gesamtleistung.

Eine zusammenfassende Gegenüberstellung der Amplituden-, Frequenz- und Phasenmodulation wird in Tafel 6.1 gezeigt. Bei der Pulsmodulation wird nicht wie bei der AM, FM und PM eine harmonische, hochfrequente Trägerschwingung verwendet, sondern eine Impulsfolge. Diese wird durch die Sendefunktion verändert. Die verschiedenen Pulsmodulationsverfahren sind digitale Modulationsverfahren. Ihr großer Vorteil gegenüber AM und WM liegt in der störungsfreien Nachrichtenübertragung.

6.2.

Modulations- und Demodulationsverfahren und -schaltungen

6.2.1.

Amplitudenmodulation und -demodulation

Modulation

Voraussetzung für eine Amplitudenmodulation ist das Vorhandensein eines Bauelements mit nichtlinearer Strom-Spannungs-Kennlinie, z. B. Diode, Transistor.

Das Grundprinzip ist dabei so, daß die durch die Signalfunktion überlagerte Trägerschwingung dem Bauelement mit nichtlinearer Kennlinie zugeführt und damit amplitudenmäßig im Rhythmus der Signalfunktion beeinflusst wird. Im Bild 6.12 wird dieser Vorgang an einer quadratischen Strom-Spannungs-Kennlinie gezeigt. Führt man den im Bild 6.12 angegebenen amplitudenmodulierten Ausgangsstrom über einen

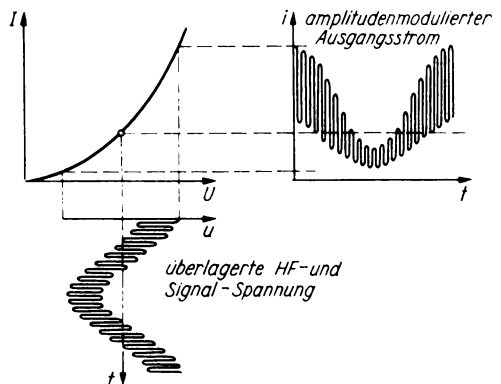


Bild 6.12. Amplitudenmodulation durch nichtlineare Strom-Spannungs-Kennlinie

Tafel 6.1. Vergleich der Modulationsarten AM, FM und PM

	AM	FM	PM
Signal- oder Sendeschwingung	$s = \hat{U}_n \cos \omega_n t$ = zeitabhängige Spannung	$s = \Delta \Omega \cos \omega_n t$ = zeitabhängige Frequenz	$s = \Delta \Phi \cos \omega_n t$ = zeitabhängiger Phasenwinkel
Unmodulierte Trägerschwingung	$u_h = [\hat{U}_h] \cos \omega_h t$	$u_h = \hat{U}_h \cos[\omega_h t]$	$u_h = \hat{U}_h \cos[\omega_h t]$
Modulierte Schwingung	$u_{AM} = [\hat{U}_h + \hat{U}_n \cos \omega_n t] \cos \omega_h t$ = $\hat{U}_h [\cos \omega_h t + \frac{1}{2} m \cos(\omega_h + \omega_n) t + \frac{1}{2} m \cos(\omega_h - \omega_n) t]$	$u_{FM} = \hat{U}_h \cos[\omega_h t + \Delta \Omega \cos \omega_n t]$ = $\hat{U}_h \cos[\omega_h t + \frac{\Delta \Omega}{\omega_n} \sin \omega_n t]$	$u_{PM} = \hat{U}_h \cos[\omega_h t + \Delta \Phi \cos \omega_n t]$
Modulationsintensität	Modulationsgrad $m = \frac{\hat{U}_n}{\hat{U}_h}$	Modulationsindex $x = \frac{\Delta \Omega}{\omega_n}$	Phasenhub $\Delta \Phi$
Bei gleichbleibender Aussteuerung konstant bleibende maximale Änderung der modulierten Größe	Amplitudenhub $\hat{U}_n$	Frequenzhub $\Delta \Omega = \Delta \Phi \omega_n$	Phasenhub $\Delta \Phi = \frac{\Delta \Omega}{\omega_n}$
Bandbreite des Frequenzspektrums	$b = 2\omega_n$	$b = 2[\Delta \Omega + (1 \dots 2)\omega_n]$	$b = 2[\Delta \Phi \omega_n + (1 \dots 2)\omega_n]$
Gesamtleistung einer modulierten Schwingung	$P_{AM} = P_h + 2P_{(h \pm n)} = P_h \left(1 + \frac{m^2}{2}\right)$	$P_{FM, PM} = P_h = \text{konst.}; P_{FM, PM} = P_0 + 2P_{(h \pm n)} = P_h \left(J_0^2 + 2 \sum_{p=1}^n J_p^2\right)$	
Trägerleistung ohne Modulation	$P_h = \text{Konst.} \cdot \hat{U}_h^2 = \frac{\hat{U}_h^2}{2R}$	$P_h = \text{konst.} \cdot \hat{U}_h^2 = \frac{\hat{U}_h^2}{2R}$	$P_h = \text{konst.} \cdot \hat{U}_h^2 = \frac{\hat{U}_h^2}{2R}$
Trägerleistung mit Modulation		$P_0 = P_h J_0^2$	
Leistung einer Seitenfrequenzschwingung	$P_{(h-n)} = \text{Konst.} \cdot U_h^2 \frac{m^2}{4} = P_h \frac{m^2}{4}$	$P_{(h \pm n)} = P_h \sum_{p=1}^n J_p^2$	

$J_0, J_1, J_2, \dots, J_n$ sind die von $\Delta \Phi$ abhängigen Bessel-Funktionswerte

auf die Trägerfrequenz abgestimmten Schwingkreis, dann fällt ohne die überlagerten Gleichanteile eine amplitudenmodulierte Spannung u_{AM} gemäß dem Bild 6.1 ab.

Diodenschaltungen

Amplitudenmodulationsschaltungen mit Dioden sind besonders einfach. Nachteilig ist, daß Dioden lediglich nichtlineare Widerstände sind und keine Verstärkungseigenschaften aufweisen. Sie werden nur bei kleinen Leistungen angewendet.

Diodenmodulator mit Stromflußwinkelmodulation

Bei einem einfachen Diodenmodulator nach Bild 6.13 wird über eine in Sperrichtung vorgespannte Diode der Stromflußwinkel der Trägerschwingung durch die Signalfrequenz gesteuert. Es können relativ große Träger- und Signalamplituden verarbeitet werden. Bei hoher Modulationsintensität steigt allerdings der Klirrfaktor stark an.

Verbesserungen gegenüber dem einfachen Diodenmodulator liefern die Gegentaktschaltung zweier Diodenmodulatoren und der Ringmodulator. Diese kompensieren am Ausgang die Trägerschwingung und einige unerwünschte Modulationsprodukte und sind damit für eine Einseitenband-Amplitudenmodulation geeignet.

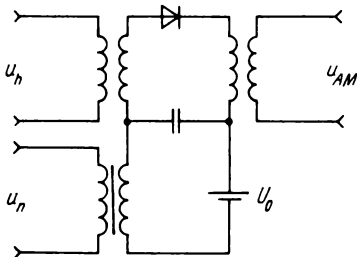


Bild 6.13. Diodenmodulation mit Stromflußwinkelmodulation

Ringmodulator

Der Ringmodulator ist eine Doppelgegentaktschaltung und hat eine große praktische Bedeutung in der Nachrichtentechnik. Bei ihm sind bei symmetrischem Aufbau alle Schwingkreise gegeneinander entkoppelt, so daß am Ausgang nur ein kleines Frequenzspektrum anliegt. Das ist vor allem deshalb erwünscht, weil sich dadurch das Nutzseitenband viel leichter von den benachbarten Frequenzen trennen läßt.

Um die Wirkungsweise zu verstehen, geht man davon aus, daß jede Diode in Sperrichtung einen großen, in Durchlaßrichtung einen kleinen Widerstand hat. Setzt man ideale Verhältnisse an, also völlige Sperrung der Dioden bzw. idealen Durchlaß, dann erhält man bei Aussteuerung durch eine Trägerschwingung mit der Frequenz ω_h als zeitlichen Widerstandsverlauf eine Umpolfunktion. Die Dioden können damit als von der Trägerschwingung gesteuerte Schalter betrachtet werden.

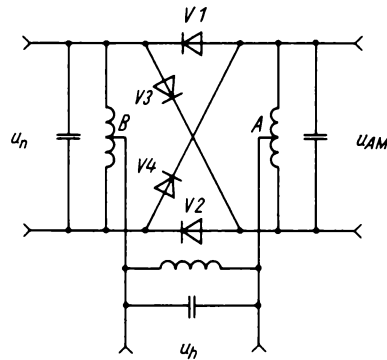


Bild 6.14. Ringmodulator

Überträgt man diese Erkenntnisse auf die Gesamtschaltung des Ringmodulators, dann ergibt sich unter Benutzung der Ersatzschaltung nach Bild 6.15a, daß während der positiven Halbwelle der Trägerschwingung (+ an A) die Dioden $V1$ und $V2$ durchlässig sowie $V3$ und $V4$ gesperrt sind. Das Modulationssignal gelangt damit direkt zum Ausgang und liefert eine Spannung, die phasengleich mit der Trägerspannung ist. Während der negativen Halbwelle der Trägerspannung (– an A) werden die Dioden $V3$ und $V4$ leitend, während $V1$ und $V2$ sperren. In diesem Fall erfährt das Modulationssignal am Ausgang eine Umpolung und liegt gegenphasig zur Trägerspannung. Dieser Vorgang wiederholt sich, so daß im Rhythmus der Trägerfrequenz eine ständige Umpolung der Modulationsspannung erfolgt. Es entsteht damit am Ausgang das im Bild 6.15b dargestellte Schwingungsbild.

Da sich beim Ringmodulator in jedem Zeitmoment die magnetischen Felder der beiden von u_h angetriebenen Teilströme aufheben, kann im Ausgangsfrequenzspektrum keine Trägerschwingung und keine ihrer Oberwellen auftreten. Da bei völliger Symmetrie die Amplituden der Sendeschwingungen beim Umpolen auf

jedem der zwei Strompfade gleich sind, verschwindet auch die ω_n -Komponente am Ausgang.

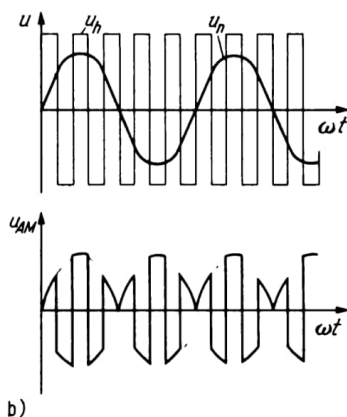
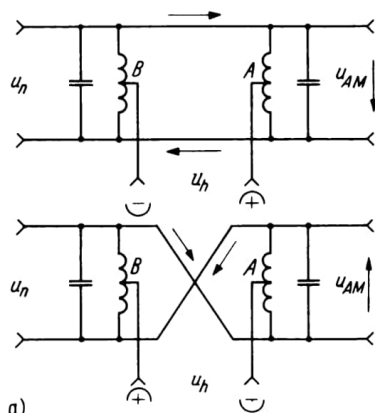


Bild 6.15. Zur Wirkungsweise des Ringmodulators

a) Ersatzschaltung; b) Schwingungsbild

Transistorschaltungen

Amplitudenmodulationsschaltungen mit Transistoren sind für kleine Leistungen geeignet und finden vor allem bei tragbaren Funksprechgeräten, Fernsteuersendern u. ä. Anwendung.

Transistoreingangs- und -ausgangsmodulator Grundsätzlich unterscheidet man zwischen Eingangs- und Ausgangsmodulation. Eine Eingangsmodulation liegt dann vor, wenn die Träger- und Signalschwingung zwischen Basis und Emitter des Transistors gelegt wird. Da diese Elektrodenanordnung eine Diode ist, kann man jede Eingangsmodulation mit Transistoren auf eine Diodenmodulation mit nachfolgender Verstärkung zurückführen. Im Bild 6.16 ist eine entsprechende Schaltung dargestellt. Während

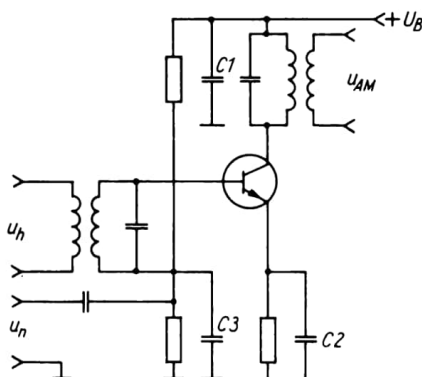


Bild 6.16. Transistoreingangsmodulator

$C1$ und $C2$ möglichst jede Wechselgröße kurzschließen sollen, darf $C3$ das nur für die Trägerfrequenzschwingung.

Bei genügend kleiner Trägerschwingung kann damit relativ verzerrungsarm moduliert werden; liegen größere Trägeramplituden vor, muß die Basis-Emitter-Diode in Sperrichtung vorgespannt werden, und es kommt zu einer Stromflußwinkelmodulation wie bei der Schaltung gemäß Bild 6.13.

Bei einer Ausgangsmodulation wird die Trägerschwingung dem Eingang und die Signalfrequenz dem Ausgang, d. h. dem Kollektor, zugeführt. Derartige Schaltungen arbeiten im Gegensatz zu denen mit Eingangsmodulation auch dann noch zufriedenstellend, wenn relativ große Trägeramplituden anliegen. Die Modulationsverzerrungen bleiben selbst bei hohem Modulationsgrad verhältnismäßig gering.

Integrierter Doppelgegentaktmodulator

Bild 6.17 zeigt die Prinzipschaltung eines Doppelgegentaktmodulators mit integriertem Schaltkreis. Die Wirkungsweise entspricht der

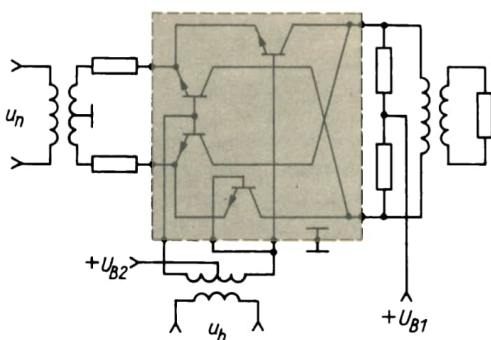


Bild 6.17. Modulator mit integriertem Schaltkreis

des bereits behandelten Ringmodulators, nur daß anstelle der Dioden npn-Anordnungen auftreten. Die sich ergebenden Vorteile sind ein einfacherer Schaltungsaufbau und eine sehr kleine Dämpfung. Die praktisch realisierten integrierten Modulatorschaltkreise haben allerdings noch weitere Schaltungsstrukturen als die im Bild 6.17 angegeben.

Demodulation

Die Demodulation ist die Umkehrung der Modulation. Durch sie wird das der hochfrequenten Trägerschwingung aufgeprägte Signal zurückgewonnen. Grundsätzlich kann jedes Bauelement mit einer nichtlinearen Strom-Spannungs-Kennlinie zur Demodulation verwendet werden. Im Interesse geringer Verzerrungen sollte diese Kennlinie eine ideale Knickkennlinie sein.

Diodendemodulator

Am gebräuchlichsten ist die Diodendemodulation. Bei ihr wird stets eine Spitzengleichrichtung angewendet, indem durch einen Kondensator C die Diode mit der Amplitude der modulierten Trägerspannung vorgespannt wird. Dadurch wird der Arbeitspunkt gerade so weit verschoben, daß – theoretisch – kein Diodenstrom mehr fließt. Da die Amplitudenänderung der modulierten Trägerfrequenzschwingung aber der Signalfrequenzschwingung entspricht, ändert sich die Diodenvorspannung mit der

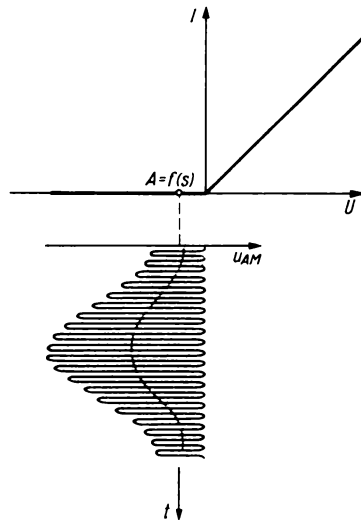


Bild 6.18. Demodulation einer AM-Schwingung durch ein Bauelement mit Knickkennlinie

Signalfunktion, und der Arbeitspunkt schwankt im Rhythmus des zu demodulierenden Signals. An einem Belastungswiderstand R_L kann diese signalabhängige Diodenvorspannung abgegriffen werden. Im Bild 6.18 ist dieser Vorgang veranschaulicht.

Es gibt zwei prinzipielle Möglichkeiten zur Schaltung von Diode und Belastungswiderstand R_L : die Reihen- und die Parallelschaltung. Zu beiden hat man die durch C und R_L gebildete Zeitkonstante so zu wählen, daß die an C und R_L abfallende Diodenvorspannung der Modulationsfrequenz, nicht aber der wesentlich höherliegenden Trägerfrequenz folgen kann, d. h.

$$T_n = \frac{1}{f_n} > \tau = CR_L > T_h = \frac{1}{f_h}$$

Zur praktischen Dimensionierung gilt die Bedingung

$$\frac{1}{\omega_{n\max}} > \tau = CR_L > \frac{10}{\omega_h} \quad (6.24)$$

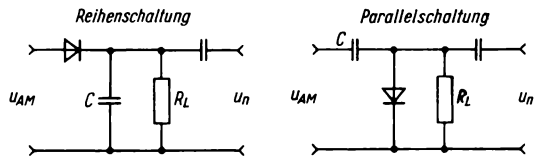


Bild 6.19. Diodendemodulationsschaltungen

Im Bild 6.19 werden die beiden Demodulationsschaltungen gezeigt. Bezüglich der Belastung des vor einer Demodulationsschaltung liegenden Schwingkreises unterscheiden sich die zwei Schaltungen. Nimmt man an, daß die Belastung durch einen gleichwertigen Dämpfungswiderstand R_d ersetzt werden kann, dann ist die darin verbrauchte Wechselleistung

$$P_{h\sim} = \frac{U_{h\text{eff}}^2}{R_d}$$

und mit

$$U_{h\text{eff}} = \frac{\hat{U}_h}{\sqrt{2}} \text{ folgt } P_{h\sim} = \frac{\hat{U}_h^2}{2 R_d}$$

Diese Leistung tritt nach der Demodulation aber auch als (Gleich-)Leistung P_L im Arbeitswiderstand R_L auf, mit einer (Gleich-)Spannung, die bei der Reihenschaltung wegen der Spitzengleichrichtung gleich dem Scheitelwert

$\hat{U}_h$ der Trägerspannung ist. Es ergibt sich dann

$$P_{h-} = P_{-}$$

$$\frac{\hat{U}_h^2}{2 R_d} = \frac{\hat{U}_L^2}{R_L}$$

$$R_{d(\text{Reihe})} = \frac{R_L}{2} \quad (6.25)$$

Bei der Parallelschaltung wirkt der Arbeitswiderstand R_L auch bei fortgenommener Diode energieentziehend, so daß dieser noch parallel zu dem in Gl. (6.25) bestimmten Dämpfungswiderstand liegt. Es ergibt sich darum

$$R_{d(\text{parallel})} = R_{d(\text{Reihe})} \parallel R_L \\ = \frac{R_L}{2} \parallel R_L$$

$$R_{d(\text{parallel})} = \frac{R_L}{3} \quad (6.26)$$

Die Belastung eines vor der Diodenmodulationsschaltung liegenden Schwingkreises wirkt so, als ob ein Dämpfungswiderstand R_d parallel zum Kreis liegt. Dabei ist bei der Reihenschaltung $R_d = R_L/2$ und der Parallelschaltung $R_d = R_L/3$. Die Diodenparallelschaltung verhält sich also dämpfungsmäßig ungünstiger als die entsprechende Reihenschaltung.

Synchron- oder Produktdemodulator

Amplitudenmodulierte Schwingungen mit relativ niedriger Trägerfrequenz und kleiner Bandbreite erfordern nur eine verhältnismäßig geringe Linearität des Demodulators, die in ausreichendem Maße vom Diodendemodulator erfüllt wird.

Bei der Demodulation des zwischenfrequenten AM-Bildsignals beim Fernsehen sind die Linearitätsanforderungen an den Demodulator wesentlich höher, da neben den benötigten Helligkeits- und Synchronsignalen (BAS) auch der Tonträger bei 5,5 MHz und der Farbträger bei 4,43 MHz entstehen. Ton- und Farbträger können bei ungenügender Linearität des Demodulators Mischprodukte bilden, die sich als unangenehme Bildstörungen (Moiré) auswirken.

Ein allen Anforderungen genügender Demodulator erfordert einen hohen Aufwand, läßt sich aber durch eine integrierte Schaltung problemlos gestalten. Neuere Fernsehempfänger arbeiten deshalb ausschließlich mit integrierten Synchrondemodulatoren.

Im Synchron- oder Produktdemodulator (im

folgenden wird nur noch der erstgenannte Begriff gebraucht) findet ein synchroner Modulationsvorgang statt, indem das Zweiseitenband- oder das Einseitenbandsignal mit einer zugeetzten Trägerschwingung moduliert wird, deren Frequenz und Phasenlage der ursprünglichen Trägerschwingung entspricht.

Wirkungsweise des Synchrondemodulators (Bild 6.20):

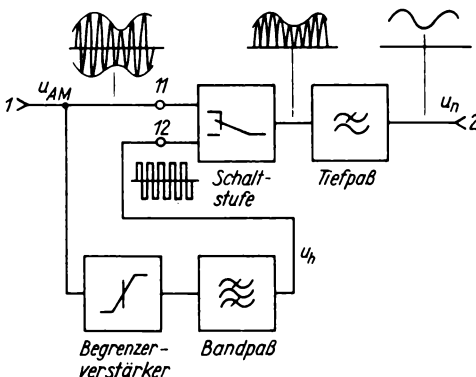


Bild 6.20. Übersichtsschaltplan eines Synchrondemodulators

Die zu demodulierende AM-Spannung u_{AM} wird gleichzeitig an den Eingang des Begrenzerverstärkers und an einen Eingang der Schaltstufe gelegt. Im Begrenzerverstärker wird die Spannung so hoch verstärkt und begrenzt, daß eine Trägerfrequenzimpulsfolge (Rechteckkurve) entsteht. Die dabei entstehenden Oberwellen werden durch den Bandpaß unterdrückt, so daß an dessen Ausgang die in der Amplitude begrenzte Trägerspannung u_h bereit steht. Die der Schaltstufe zugeführte AM-Spannung wird nun durch die über dem zweiten Eingang angelegte Trägerspannung im Rhythmus der Trägerfrequenz ω_h durchgeschaltet. Im Ausgangskreis der Schaltstufe führt das zu einer dem momentanen Modulationsgrad entsprechenden Stromänderung, die der Modulationsspannung u_n proportional ist. Unerwünscht auftretende Schwingungsanteile, besonders die mit der doppelten Trägerfrequenz, werden durch den nachgeschalteten Tiefpaß beseitigt.

Kernstück des Synchrondemodulators ist die im folgenden etwas näher betrachtete Schaltstufe (Bild 6.21). Sie besteht aus drei Differenzverstärkerstufen $V1/V2$, $V3/V4$, $V5/V6$ und der Stromquelle $V7$. Der Verstärker $V1/V2$ hat durch die beiden Emitterwiderstände einen großen linearen Aussteuerungsbereich. Die beiden

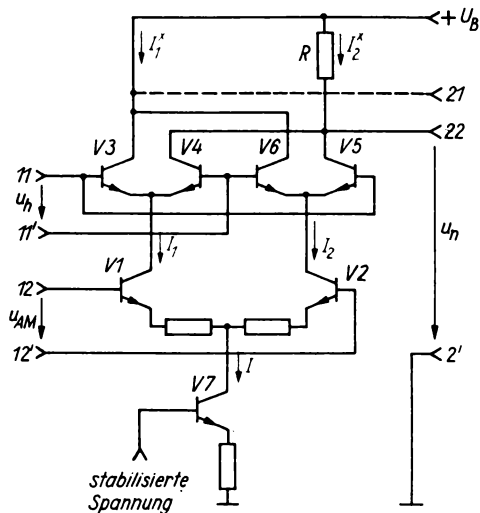


Bild 6.21. Schaltstufe eines Synchrondemodulators (Doppelgegentaktmischer)

oberen kreuzgekoppelten Differenzverstärker entsprechen der Schaltung im Bild 6.17.

Die Schaltung arbeitet so, daß sich der von der Stromquelle angetriebene Strom I bei fehlender Ansteuerung gleichmäßig auf die Transistoren $V1$ und $V2$ sowie weiter auf $V3$ und $V4$ bzw. $V5$ und $V6$ aufteilt. Die Ruhestrome I_1^* und I_2^* haben also eine Stärke von je zwei Viertel des Gesamtstroms:

$$I_1^* = I_2^* = \frac{2}{4} I = \frac{1}{2} I.$$

Werden der Differenzverstärker $V1/V2$ oder der kreuzgekoppelte Differenzverstärker getrennt angesteuert, ändern sich die Ruhestrome nicht und behalten ihren Wert $I_1^* = I_2^* = \frac{1}{2} I$. Liegt

beispielsweise an $V1/V2$ eine Steuerspannung, die den Transistor $V2$ sperrt und den Strom nur durch $V1$ fließen läßt, dann wird dieser durch die Transistoren $V3$ und $V4$ halbiert, und es fließen die gleichen Ruhestrome wie im oben beschriebenen Fall ohne Aussteuerung. Zum gleichen Ergebnis kommt man für andere Aussteuerungsverhältnisse – gleichgültig, ob sie an $V1/V2$ oder an $V3$ bis $V6$ auftreten. Erst wenn gleichzeitig an beiden Eingängen Steuerspannungen anliegen, ändern die Ströme I_1^* und I_2^* ihre Größe durch den sie überlagernden u_{AM} -Verlauf.

Die im Bild 6.22 dargestellten Zeitverläufe zeigen die beschriebenen Vorgänge. An die Ein-

gänge 11–11' wird die vom Begrenzerverstärker gelieferte Trägerspannung u_h gelegt. Die Amplitude des Trägers ist dabei so groß, daß die Transistoren in Schalterbetrieb arbeiten. Die positiven Amplituden öffnen die Transistoren $V3$ und $V5$, während sie $V4$ und $V6$ sperren; bei den negativen Amplituden ist es umgekehrt. Die an 12–12' anliegende AM-Spannung u_{AM} wird dadurch im Rhythmus der Trägerfrequenz ω_h voll übertragen bzw. unterbrochen. Der Ruhestrom I_1^* wird damit um einen der AM-Spannung entsprechenden Wert vergrößert, I_2^* dagegen vermindert. Am Widerstand R verursacht I_2^* einen Spannungsabfall, dessen Mittelwert die Modulationsspannung u_n ist und am Anschluß 22 zur Verfügung steht. Die an R abfal-

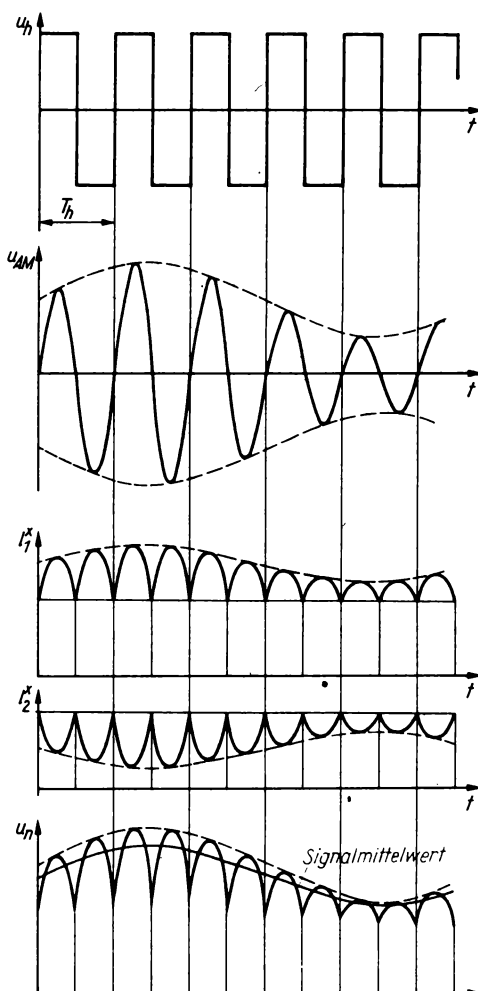


Bild 6.22. Zur Wirkungsweise des Synchrondemodulators

lende Spannung ist eine Schwingung, die außer der gewünschten Frequenz ω_n noch einen Gleichanteil und die doppelte Trägerfrequenz $2\omega_n$ enthält. Der meist nachfolgende Videovorverstärker übernimmt die Funktion des Tiefpasses und unterdrückt die unerwünschten Schwingungen. Auf die Auskoppelmöglichkeit über einen Anschluß 21 wird gewöhnlich verzichtet.

AM-Synchrodemodulatoren für Bildsignale werden meist gemeinsam mit einem regelbaren ZF-Verstärker, einem Videovorverstärker, einer Taststufe, einem Regelverstärker und einem Schwellwertverstärker für die Tunerregelung kombiniert integriert. Der VEB Halbleiterwerk Frankfurt (Oder) fertigt einen derartigen integrierten Schaltkreis mit der Typenbezeichnung A 240 D.

6.2.2.

Winkelmodulation und -demodulation

Modulation

Im Abschn. 6.1.2. wurde auf die enge Verwandtschaft von Frequenz- und Phasenmodulation hingewiesen. Sie ergab sich, da die Frequenz der Geschwindigkeit der Phasenänderung entspricht. Über die Beziehung

$$\frac{\Delta\Omega}{\omega_n} = \Delta\Phi$$

lassen sich beide Modulationsarten ineinander überführen und mathematisch einheitlich behandeln.

Man gewinnt folglich aus einem phasenmodulierten Sender eine frequenzmodulierte Schwingung, wenn das niederfrequente Signal vor der Modulationsstufe ein Integrierglied durchläuft. Umgekehrt erhält man von einem frequenzmodulierten Sender eine phasenmodulierte Schwingung, wenn die modulierende Größe durch ein davorliegendes Differenzierglied geschickt wurde.

Es handelt sich also um *indirekte* Modulationsverfahren, da die FM-Schwingungen über einen phasenmodulierten Sender bzw. die PM-Schwingungen über einen frequenzmodulierten Sender gewonnen werden.

Die im folgenden behandelten Schaltungen arbeiten nach direkten Modulationsverfahren.

Frequenzmodulator

Die Prinzipschaltung einer Modulationsstufe für direkte Frequenzmodulation wird im Bild 6.23 gezeigt.

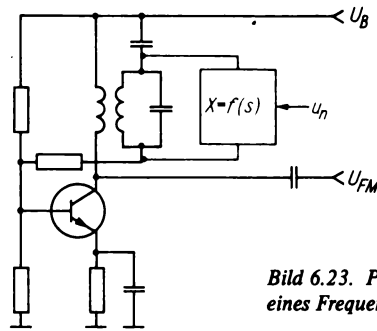


Bild 6.23. Prinzipschaltung eines Frequenzmodulators

Dem frequenzbestimmenden Schwingkreis des Oszillators wird ein durch die Modulationsspannung steuerbarer Blindwiderstand X parallelgeschaltet. Damit ändert sich die Oszillatorfrequenz ω_n als Funktion der Signalspannung, und man erhält eine frequenzmodulierte Ausgangsspannung mit der Zeitfunktion nach Gl. (6.14). Steuerbare Blindwiderstände gibt es viele. Als besonders wichtig sind die mit Transistoren zu realisierenden Reaktanzschaltungen. Eine von den vier möglichen Schaltungsvarianten wird im folgenden besprochen.

Im Bild 6.24 wird eine Reaktanzschaltung gezeigt, die wie ein induktiver Blindwiderstand wirkt.

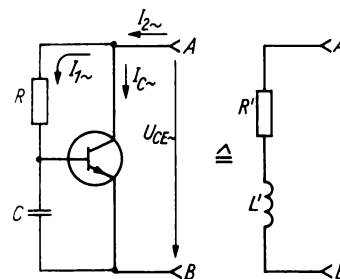


Bild 6.24. Induktive Reaktanzschaltung

Eine der Gleichspannung überlagerte Wechselspannung U_{CE} wird außen zwischen A und B angelegt. Sie treibt einen Wechselstrom I_2 an, der sich in zwei Teilströme I_1 und I_C aufteilt. Der viel kleinere Strom I_1 fließt durch die parallel zum Transistor liegende Reihenschaltung R und C . Wird R viel größer als $X_C = \frac{1}{\omega C}$

gewählt, ist I_1 nahezu phasengleich mit U_{CE} . I_1 verursacht an C einen um 90° nacheilenden Spannungsabfall U_C . Da dieser zwischen Basis und Emitter des Transistors liegt, steuert er den mit U_C in Phase liegenden Kollektorwechsel-

strom $I_{C\sim}$. Wegen des im Verhältnis zu $I_{C\sim}$ viel kleineren Stromes $I_{1\sim}$ gilt in erster Näherung $I_{C\sim} \approx I_{2\sim}$. $I_{2\sim}$ eilt damit der Spannung $U_{CE\sim}$ fast um 90° nach. Ein solches Verhalten zeigt aber auch eine verlustbehaftete Spule mit dem Widerstand R' und der Induktivität L' . Das Zeigerdiagramm (Bild 6.25) zeigt die erläuterten Zusammenhänge. Wegen der nichtlinearen Strom-Spannungs-Kennlinie des Transistors ist aber der Kollektorwechselstrom $I_{C\sim} \approx I_{2\sim}$ arbeitspunktabhängig. Wird nun die Basis-Emitter-Spannung im Rhythmus der Modulationsspannung geändert, ändert sich auch X_{ind} .

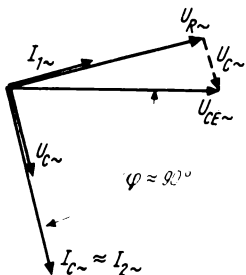


Bild 6.25. Zeigerdiagramm einer induktiv wirkenden Reaktanzschaltung

Weitere Schaltungsvarianten erhält man, wenn die Schaltelemente R und C vertauscht werden oder anstelle des Kondensators C eine Spule L eingefügt wird.

Phasenmodulator

Von großer praktischer Bedeutung ist der im Bild 6.26 angegebene Filterphasenmodulator. Mittels des durch die Modulationsspannung steuerbaren Blindwiderstands X wird der am Ausgang des Filters liegende Schwingkreis verstimmt. Damit ändern sich im Rhythmus der Signalfrequenz die Amplitude und der Phasenwinkel der Ausgangsspannung.

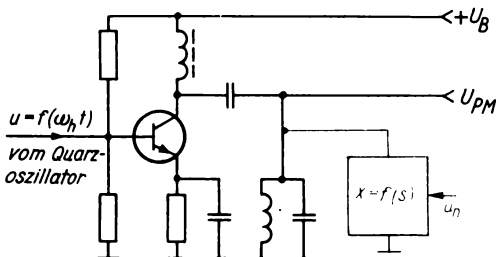


Bild 6.26. Prinzipschaltung des Filterphasenmodulators

Während die Amplitudenänderung unerwünscht ist und durch eine Begrenzerschaltung beseitigt werden muß, ist die Phasenwinkeländerung die erwartete Phasenmodulation. Man erhält somit eine phasenmodulierte Ausgangsspannung mit der Zeitfunktion nach Gl. (6.17). Die Modulationskennlinie des Filterphasenmodulators entspricht dem im Bild 6.27 angegebenen Phasenverlauf eines Parallelschwingkreises.

Als steuerbare Blindwiderstände können die bereits behandelten Reaktanzschaltungen verwendet werden. Der ausnutzbare maximale Phasenhub beträgt beim Filterphasenmodulator mit Einzelkreis und 1 % Klirrfaktor nur $\Delta\Phi \approx 20^\circ \approx 0,35$.

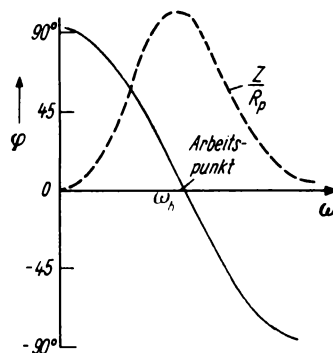


Bild 6.27. Prinzipieller Verlauf des Phasenwinkels und Scheinwiderstands eines Parallelschwingkreises als Funktion der Frequenz

Demodulation

Bei der Demodulation winkelmodulierter Schwingungen könnte man streng zwischen Frequenz- und Phasenmodulation unterscheiden. Das wäre aber formal, da in der Praxis meistens ein Kompromiß zwischen FM und PM gemacht wird (s. Bild 6.3). Man nennt deshalb (nicht ganz exakt) alle derartigen Demodulatoren *FM-Demodulatoren*.

Es gibt vier unterschiedliche Verfahren, um frequenz- oder phasenmodulierte Schwingungen zu demodulieren:

- FM-AM-Wandlung, Rückgewinnung der Sendeschwingung durch AM-Demodulation Schaltungen: Flankendemodulator, Phasendiskriminator (Riegger-Kreis), Verhältnisdiskriminator (Ratiodetektor).
- FM-PDM-Wandlung, Rückgewinnung der Sendeschwingung durch einen Tiefpaß

Schaltung: Koinzidenzdemodulator (Phasendemodulator, Quadraturdemodulator).

- FM-PFM-Wandlung, Rückgewinnung der Sendeschwingung durch einen Tiefpaß
Schaltung: Zählerdiskriminator (Nulldurchgangsdiskriminator).
- Phasenregelung
Schaltung: PLL-Demodulator (Phasenregelkreis, Phasensynchronfilter, Nachlauffilter).

Demodulator mit FM-AM-Wandlung

Den prinzipiellen Aufbau eines Demodulators mit FM-AM-Wandlung zeigt Bild 6.28.

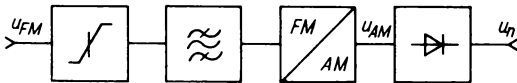


Bild 6.28. Prinzip eines FM-Demodulators mit FM-AM-Wandlung

Die frequenz- oder phasenmodulierte Schwingung wird in einem Begrenzer von unerwünschten Amplitudenänderungen, die unmittelbar in die Amplitude des demodulierten Signals eingehen würden, befreit. Der sich anschließende Bandpaß sperrt die durch Kurvenverformungen entstandenen Oberwellen des Trägers. Diese Aufgabe kann auch vom Demodulationsfilter, das die FM- in eine AM-Schwingung umwandelt, mit übernommen werden. Dem FM-AM-Wandler folgt ein AM-Demodulator, durch den die niederfrequente Signalspannung gewonnen wird. Die Schaltungseinheit mit FM-AM-Wandler und AM-Demodulator wird *Diskriminator* genannt.

Im einfachsten Fall kann als Demodulationsfilter ein Schwingkreis benutzt werden. Der Arbeitspunkt muß dabei auf der Flanke des Schwingkreises, möglichst im Wendepunkt der Schwingkreiskennlinie liegen. Da der maximal zulässige Hub durch den Verlauf der Flanke begrenzt und die Linearität der Kennlinie ungenügend ist, kommt nur eine Gegentaktschaltung in Frage.

Der *Gegentaktsflankendemodulator* hat zwei um $\pm \Delta\omega$ gegenüber der Trägerfrequenz verstimmte Schwingkreise. Nach Gleichrichtung können die gegeneinandergeschalteten niederfrequenten Spannungen u_{n1} und u_{n2} als Differenz entnommen werden, und man erhält als gewünschte Signalspannung

$$u_n = u_{n1} - u_{n2}.$$

Im Bild 6.29 sind die Schaltung und die dazugehörige Demodulationskennlinie angegeben. Durch verschiedene Schaltungsmaßnahmen kann die Demodulationskennlinie weiter linearisiert werden.

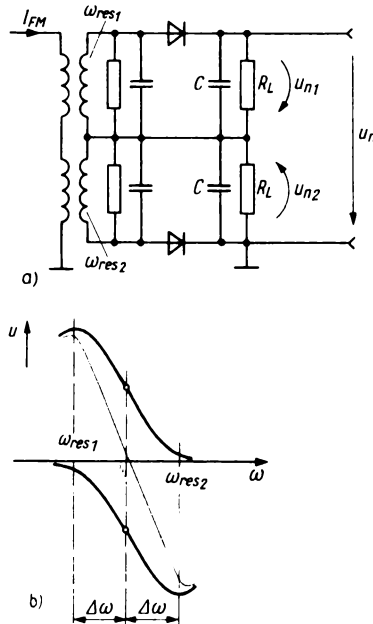


Bild 6.29. Gegentaktsflankendemodulator

a) Schaltung; b) Demodulationskennlinie

Beim Phasendiskriminator (Riegger-Kreis) ersetzt man die zwei gegeneinander verstimmten Schwingkreise durch einen einzigen mit Mittelanzapfung. In die Mittelanzapfung des Sekundärkreises wird zusätzlich durch den Kondensator C_k die Spannung des Primärkreises eingekoppelt. Über die Drosselspule L wird der Gleichstromkreis geschlossen. Im Bild 6.30 ist die Schaltung dieses wichtigen Demodulators dargestellt.

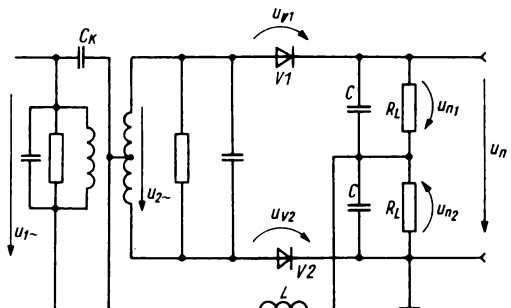


Bild 6.30. Phasendiskriminator (Riegger-Kreis)

Da über die Kondensatoren C die Katoden der Dioden hochfrequenzmäßig auf Massepotential liegen, tritt über jeder Diode die geometrische Summe der Primär- und halben Sekundärspannung auf, also

$$u_{v1} = u_{1\sim} + \frac{u_{2\sim}}{2} \quad \text{und} \quad u_{v2} = u_{1\sim} - \frac{u_{2\sim}}{2}.$$

Über den Widerständen R_L fallen folglich die Beträge der Diodenspannungen ab, d. h.

$$u_{a1} = |u_{v1}|$$

$$u_{a2} = |u_{v2}|.$$

Wegen der Gegeneinanderschaltung liegt dann am Ausgang die Signalspannung als Differenz der beiden, also

$$u_n = u_{a1} - u_{a2}.$$

Wie aus der Theorie hervorgeht, ist beim Bandfilter im Resonanzfall die Sekundärspannung gegenüber der Primärspannung um 90° phasenverschoben; der Phasenwinkel beträgt damit $\varphi = 90^\circ$.

Stimmt bei vorliegender Schaltung nach Bild 6.30 die Augenblicksfrequenz genau mit der Mittenfrequenz $\omega_h = \omega_{res}$ überein (unmodulierter Träger), dann sind die Zeigerlängen der beiden Diodenspannungen, d. h. ihre Beträge, gleich groß, und die Ausgangssignalspannung ist $u_n = 0$.

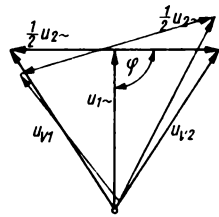


Bild 6.31. Zeigerdiagramm des Phasendiskriminators

Im Bild 6.31 wird das dazugehörige Zeigerdiagramm gezeigt. Stimmt infolge einer Modulation die Augenblicksfrequenz nicht mehr mit der Mittenfrequenz überein, dann ändert sich die Phasenverschiebung, d. h. $\varphi \neq 90^\circ$, und es entstehen die im Bild 6.31 rot eingezeichneten Verhältnisse. Da die Beträge der Diodenspannungen ungleich sind, also $|u_{v1}| \neq |u_{v2}|$, entsteht am Ausgang die der Frequenzänderung proportionale Signalspannung u_n .

Der *Verhältnisdiskriminator* (Ratiodektor) weist prinzipiell die gleiche Kreisanordnung wie der Phasendiskriminator auf, nur daß im Gegensatz zu diesem eine Diode umgepolt ist und die Ausgangsklemmen in einer Brücke liegen (Bild 6.32). Wie die Zeigerdarstellung im Bild 6.33 zeigt, ergibt sich dabei zwischen den Punkten a und b eine Differenzspannung, die in Größe und Richtung der Frequenzänderung proportional ist und damit der Signalspannung entspricht. Charakteristisch für den Verhältnisdiskriminator ist seine Begrenzerwirkung, die durch das zwischen den Punkten c und d liegende $R_L C_L$ -Glieder mit einer Zeitkonstante $\tau = 0,1 \dots 0,2$ s in Verbindung mit den beiden Dioden zustande kommt. Es kann damit oft auf eine spezielle Begrenzerstufe verzichtet werden, so daß der Verhältnisdiskriminator trotz geringerer Empfindlichkeit gegenüber dem Phasendiskriminator die bevorzugte Demodulationsschaltung

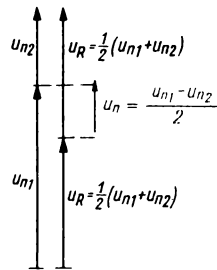
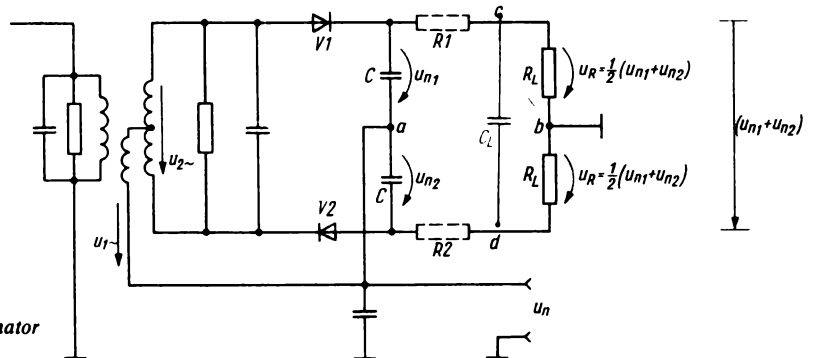


Bild 6.33. Zeigerdiagramm der Brückenschaltung des Verhältnisdiskriminators (Ratiodektor)

Bild 6.32. Verhältnisdiskriminator (Ratiodektor)



tung in den mit diskreten Bauelementen aufgebauten Empfängern der Unterhaltungselektronik ist.

Demodulator mit FM-PDM-Wandlung

Den prinzipiellen Aufbau eines Demodulators mit FM-PDM-Wandlung zeigt Bild 6.34.

Die frequenz- oder phasenmodulierte Schwingung durchläuft den Begrenzer und den Bandpaß und gelangt an den FM-PDM-Wandler. In diesem wird die Schwingung in eine Impulsfolge umgewandelt. Bei einer Änderung der Frequenz infolge der Modulation ändert sich dabei proportional die Impulsdauer. Der nachgeschaltete Tiefpaß bildet daraus den Mittelwert, der der Impulsdauer und damit der niederfrequenten Signalspannung entspricht.

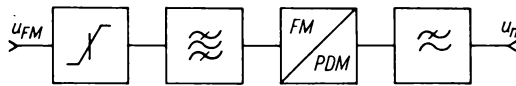


Bild 6.34. Prinzip des FM-Demodulators mit FM-PDM-Wandlung

Der FM-PDM-Wandler basiert auf folgendem Prinzip: Zwei Signale, z. B. Spannungen mit rechteckförmigem Verlauf, werden einer Koinzidenzschaltung zugeführt [coincidence, lat., (zeitliches) Zusammentreffen]. Diese liefert nur dann eine Ausgangsspannung, wenn beide Eingangsspannungen einen positiven Augenblickswert aufweisen (vgl. logisches UND-Verknüpfungsglied). Wird nun infolge der Modulation die Phase der einen Eingangsspannung gegenüber der anderen verändert, ändert sich die Dauer der Spannungsimpulse am Ausgang der Schaltung.

Die Wirkungsweise des *Koinzidenzdemodulators* (Phasen-, Quadratur- oder Vierquadrantendemodulator) ist folgende (Bild 6.35): Die zu demodulierende, stark begrenzte FM-Spannung

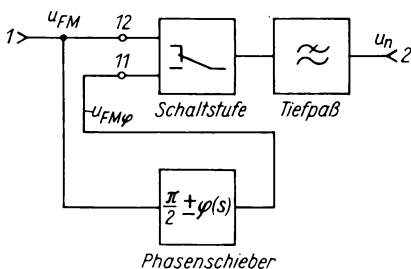


Bild 6.35. Übersichtsschaltplan eines Koinzidenzdemodulators

u_{FM} wird gleichzeitig einer Schaltstufe und einem Phasenschieber zugeführt. Bei unmodulierter Trägerspannung verursacht der Phasenschieber eine Phasenverschiebung von $\pi/2 = 90^\circ$. Diese ändert sich aber im Rhythmus der Modulation entsprechend der Phasensteilheit des Phasenschiebers.

Als Phasenschieber wäre eine Verzögerungsleitung geeignet, bei der sich die Phase proportional zur Frequenz verschiebt. Einfacher läßt sich ein Phasenschieber mit einem Schwingkreis und einem Koppelkondensator realisieren. Auch Schwingquarze oder Keramikfilter können verwendet werden.

Die im Phasenschieber aufbereitete Spannung wird der Schaltstufe zugeführt. Diese Spannung ist wegen der Selektionswirkung des Schwingkreises sinusförmig. Die Amplituden sind jedoch so groß, daß die angesteuerten Transistoren im Schalterbetrieb arbeiten. Damit ein Strom fließen kann, müssen die Transistoren in der Schaltstufe sowohl von der direkt zugeführten FM-Spannung als auch von der phasenverschobenen Spannung geöffnet sein. Infolge der Modulation ändert sich aber ihre Phasenlage zueinander, so daß sich dauermodulierte Stromimpulse konstanter Amplitude ergeben. Diese verursachen an einem Widerstand proportionale Spannungsimpulse. Im nachgeschalteten Tiefpaß wird von diesem der Mittelwert gebildet, der der niederfrequenten Signalspannung entspricht.

Bild 6.36 zeigt die vereinfachte Schaltung der Schaltstufe. Sie besteht aus zwei Differenzverstärkerstufen $V1/V4$, $V2/V3$ und der Stromquelle $V5$. Der untere Differenzverstärker wird direkt mit der begrenzten FM-Spannung angesteuert; am oberen liegt die phasenverschobene FM-Spannung. Ein Strom i_2 kann nur dann fließen, wenn gleichzeitig die Transistoren $V1$ und $V2$ geöffnet sind. Das ist nur dann der Fall, wenn die Spannung $u_{FM\varphi}$ negativ und u_{FM} positiv ist. Verschiebt sich nun infolge der Modulation die Phasenlage der beiden Spannungen zueinander, so ändert sich die gleichzeitige Öffnungsdauer der Transistoren $V1$ und $V2$ und damit die Dauer des Stromflusses durch $V2$. Als Ergebnis erhält man eine dauermodulierte Stromimpulsfolge im Ausgangskreis.

Durch i_2 werden am Widerstand R proportionale Spannungsimpulse verursacht, aus denen am Kondensator C ein Spannungsmittelwert gebildet wird. Dieser entspricht der Impulsdauer und damit der Modulationsspannung.

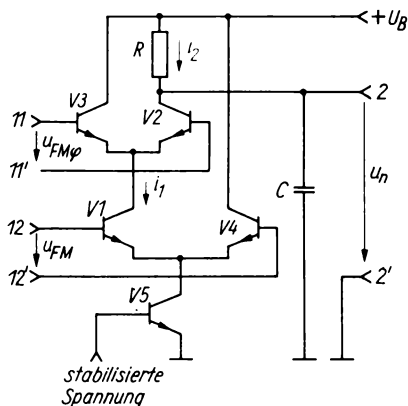


Bild 6.36. Schaltstufe des Koinzidenzdemodulators

Die im Bild 6.37 ausgewählten Zeitverläufe zeigen die beschriebenen Vorgänge. Wird z. B. infolge der Modulation die Frequenz der Trägerschwingung kleiner, so kommt es zu einer Impulsverlängerung, im umgekehrten (nicht dargestellten) Fall zu einer Impulsverkürzung.

In der Praxis werden Koinzidenzdemodulatoren mit integrierten Schaltkreisen realisiert. Die Schaltstufe ist dabei wie die des Synchronmodulators (s. Bild 6.21) aufgebaut, enthält also zwei kreuzgekoppelte Differenzverstärker. Die prinzipielle Wirkungsweise ist die entsprechend

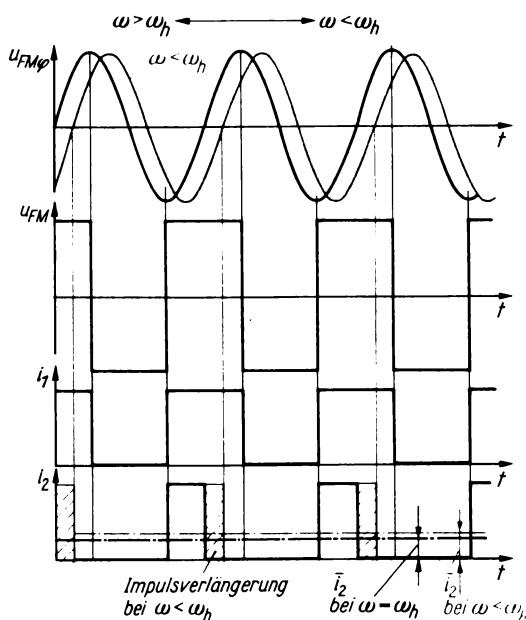


Bild 6.37. Zur Wirkungsweise des Koinzidenzdemodulators

Bild 6.36, nur mit dem Vorteil, daß sich neben der doppelten Demodulationsausgangsspannung auch eine höhere Stabilität wegen des vollkommen symmetrischen Aufbaus ergibt. Der VEB Halbleiterwerk Frankfurt (Oder) fertigt einen derartigen integrierten Schaltkreis mit der Typenbezeichnung A 220D. Dieser Schaltkreis wird vor allem als FM-ZF-Verstärker mit Demodulator im Ton-ZF-Teil von Fernsehgeräten eingesetzt, ist aber auch für andere Anwendungsgebiete geeignet. Der A 220D enthält einen 8stufigen Begrenzerverstärker, den Koinzidenzdemodulator, einen NF-Vorverstärker mit Lautstärkeregelung und eine Stabilisierungsschaltung für die Versorgungsspannung. Bild 6.38 zeigt eine typische Schaltung des A 220D als 5,5-MHz-Ton-ZF-Verstärker mit Demodulator.

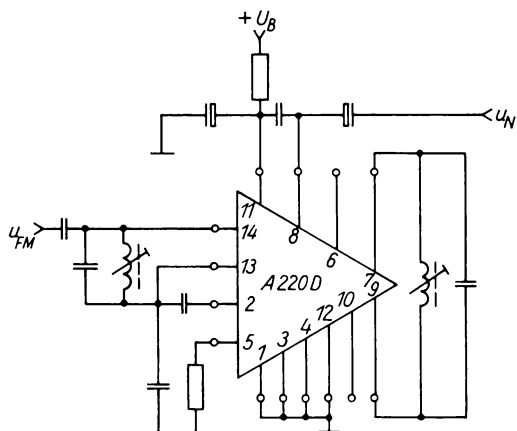


Bild 6.38. Schaltung des integrierten Schaltkreises A 220 D als Ton-ZF-Verstärker mit Demodulator

Demodulator mit FM-PFM-Wandlung

Den prinzipiellen Aufbau eines Demodulators mit FM-PFM-Wandlung zeigt Bild 6.39.

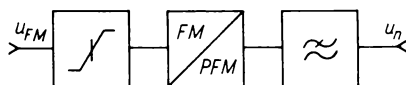


Bild 6.39. Prinzip eines FM-Demodulators mit FM-PFM-Wandlung

Die frequenz- oder phasenmodulierte Schwingung durchläuft den Begrenzer und gelangt an den FM-PFM-Wandler. In diesem wird die Schwingung in eine Impulsfolge umgewandelt. Bei einer Änderung der Frequenz infolge der

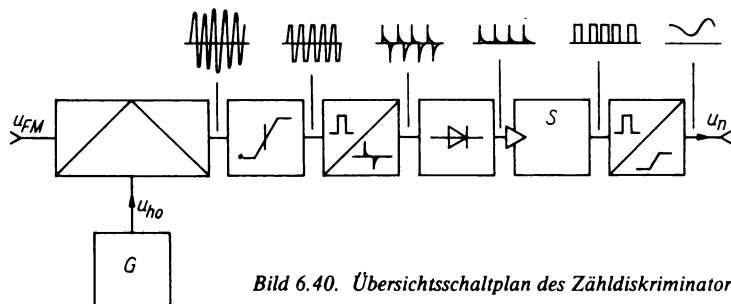


Bild 6.40. Übersichtsschaltplan des Zähldiskriminators

Modulation ändert sich dabei proportional die Impulsfrequenz. Der nachgeschaltete Tiefpaß bildet daraus den Mittelwert, der der Anzahl der Impulse je Zeiteinheit und damit der niederfrequenten Signalspannung entspricht.

Der Zähldiskriminator arbeitet nach diesem Verfahren. Da in ihm die je Zeiteinheit auftretenden Nulldurchgänge der modulierten Schwingung ausgewertet werden, wird er auch als Nulldurchgangsdiskriminator bezeichnet.

Die Wirkungsweise des Zähldiskriminators ist folgende (Bild 6.40): Die zu demodulierende FM-Spannung muß meist wegen des üblichen kleinen Hubes in einer Mischstufe auf etwa 200 kHz herabgesetzt werden. In der nachfolgenden stark wirkenden Begrenzerstufe wird jede eventuell auftretende Amplitudenänderung beseitigt und die Schwingung einem Differenzierglied zugeführt. Von der dort entstehenden Impulsfolge werden nur die positiven Nadelimpulse verwertet und der als Impulsformer arbeitenden monostabilen Kippstufe zugeführt. Diese liefert dann am Ausgang Impulse konstanter Amplitude und Dauer, deren Anzahl je Zeiteinheit aber der Momentanfrequenz proportional ist. Im anschließenden Integrierglied (Tiefpaß) wird von dieser Impulsfolge der Mittelwert gebildet, der der niederfrequenten Signalspannung entspricht.

Der Zähldiskriminator hat eine im weiten Bereich lineare Demodulationskennlinie, verzerrt also nur sehr wenig. Wegen des erforderlichen Aufwands wird er z.Z. aber noch selten in Geräten der Unterhaltungselektronik eingesetzt.

Demodulator mit Phasenregelung

Den prinzipiellen Aufbau eines Demodulators mit Phasenregelung zeigt Bild 6.41.

Einer Phasenvergleichstufe (PV), die z. B. mit dem integrierten Schaltkreis A 220D aufgebaut werden kann, wird die frequenz- oder phasen-

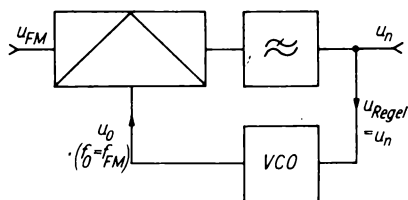


Bild 6.41. Prinzip eines FM-Demodulators mit Phasenregelung

modulierte Schwingung und die Schwingung eines spannungsgesteuerten Oszillators (VCO) zugeführt (VCO voltage controlled oscillator, engl., spannungsgesteuerter Oszillator). Eine Begrenzung der FM-Spannung ist vorteilhaft, aber nicht unbedingt erforderlich. Die Ausgangsspannung der Phasenvergleichsstufe durchläuft einen speziellen Tiefpaß, wodurch die Frequenz des VCO verändert wird. Der Oszillator wird damit phasenstarr auf die Frequenz der FM-Schwingung synchronisiert. Die Anordnung ist also ein Regelkreis (PLL phase locked loop, engl., Phasenregelkreis). Ändert sich infolge der Modulation die Eingangs- gegenüber der Oszillatorfrequenz, die zunächst auf die des unmodulierten Trägers eingestellt ist, dann entsteht hinter dem Tiefpaß eine Gleichspannung. Mit dieser wird der VCO so nachgeregelt, daß Frequenz- und Phasenabweichungen zwischen Eingangs- und Oszillatorspannung Null werden. Die Regelspannung ist also gleich der modulierten Größe und entspricht der niederfrequenten Signalspannung.

6.2.3.

Pulsmodulation und -demodulation

Modulation

Die im folgenden angegebenen Pulsmodulationsverfahren und -schaltungen sind lediglich prinzipieller Art. Sie stellen eine Auswahl verschiedener Realisierungsmöglichkeiten dar.

Pulsamplitudenmodulation PAM

Zur PAM ist ein elektronischer Schalter erforderlich, der kurzzeitig während der Dauer der Abtastimpulse geschlossen und in den Pausen geöffnet ist. Bild 6.42 zeigt eine zur PAM geeignete Diodenschaltung. Die Diode leitet, wenn die Impulsamplitude größer ist als die Summe aus der anliegenden negativen Vorspannung und dem Momentanwert der abzutastenden Sendespannung. Am Ausgang entsteht die im Bild 6.43 dargestellte amplitudenmodulierte Impulsfolge.

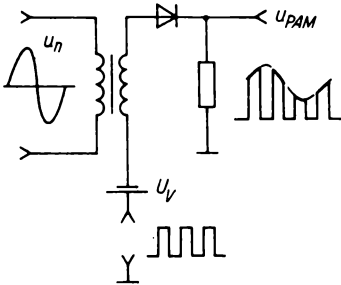


Bild 6.42. Diodenschaltung zur PAM

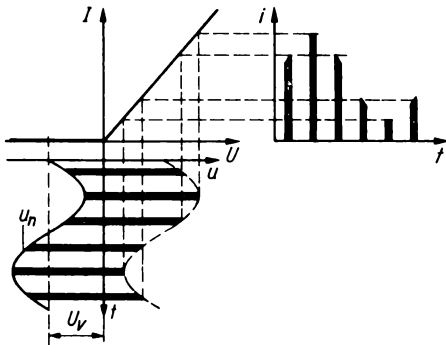


Bild 6.43. Zur Wirkungsweise der PAM-Diodenschaltung

Pulsdauermodulation PDM

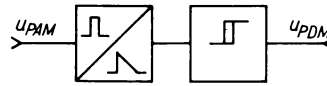


Bild 6.44. Prinzipschaltung des Pulsdauermodulators

Die PDM erfordert zunächst amplitudenmodulierte Impulse, die mit einer Schaltung nach Bild 6.42 gewonnen werden können. Aus diesen werden z. B. durch Aufladung eines Kondensators und stromkonstanter Entladung dreieckförmige amplitudenmodulierte Impulse erzeugt. Eine anschließende Schwellwertschaltung (s. Abschn. 5.4.7.) wandelt diese in dauermodulierte Impulse mit konstanter Amplitude um. Die Bilder 6.44 und 6.45 zeigen die prinzipielle Schaltung zur PDM und die entstehende Impulsfolge.

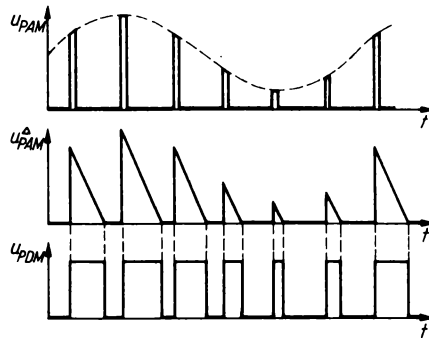


Bild 6.45. Zur Wirkungsweise der PDM-Schaltung

Pulsphasenmodulation PPM

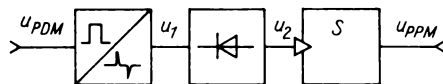


Bild 6.46. Prinzipschaltung des Pulsphasenmodulators

Die PPM kann so erfolgen, daß die mit der Schaltung nach Bild 6.44 gewonnenen dauermodulierten Impulse differenziert und gleichgerichtet werden. Die von der Rückflanke der Rechteckimpulse verursachten Nadelimpulse steuern eine monostabile Kippstufe (s. Abschnitt 5.4.6.), an deren Ausgang daraufhin eine phasenmodulierte Rechteckimpulsfolge entsteht. Die Bilder 6.46 und 6.47 zeigen die prinzipielle Schaltung zur PPM und die entstehende Impulsfolge.

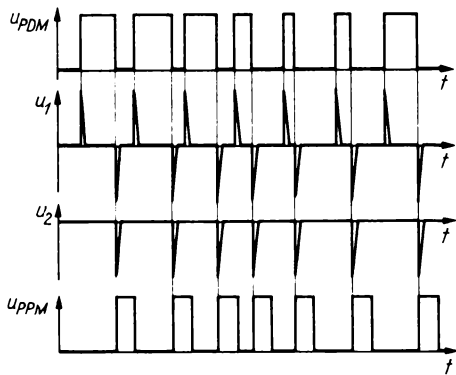


Bild 6.47. Zur Wirkungsweise der PPM-Schaltung

Pulsfrequenzmodulation PFM

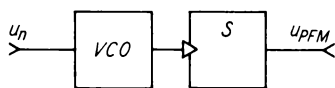


Bild 6.48. Prinzipschaltung des Pulsfrequenzmodulators

Zur PFM ist ein spannungsgesteuerter Oszillator (VCO) erforderlich, dessen Frequenz durch die Modulationsspannung verändert wird. Im einfachsten Fall könnte das ein spannungsgesteuerter astabiler Kippgenerator sein (s. Abschn. 4.3.2.). Die von diesem abgegebene Impulsfolge steuert mit den Vorderflanken die nachfolgende monostabile Kippstufe. Am Ausgang entsteht daraufhin eine Impulsfolge mit Impulsen konstanter Amplitude und konstanter Dauer, jedoch mit einer der Modulationsspannung proportionalen Impulsfrequenz. Die Bilder 6.48 und 6.49 zeigen die prinzipielle Schaltung zur PFM und die entstehende Impulsfolge.

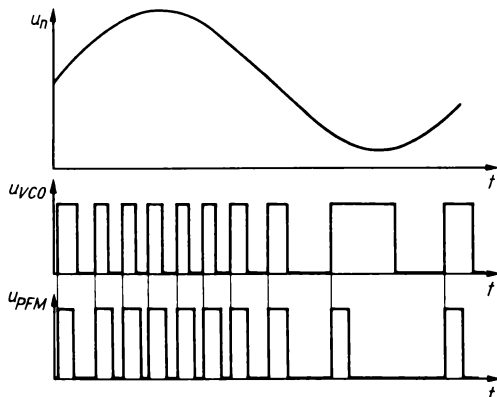


Bild 6.49. Zur Wirkungsweise der PFM-Schaltung

Demodulation

Bei amplitudenmodulierten, dauermodulierten und frequenzmodulierten Impulsfolgen entspricht der Spannungsmittelwert der Modulationsspannung. Die Signalerückgewinnung kann folglich für PAM, PDM und PFM durch einen Tiefpaß erfolgen, in dem der Spannungsmittelwert gebildet wird.

Um eine PPM-Schwingung zu demodulieren, muß diese in eine PAM, PDM oder PFM umgewandelt werden; danach wird, wie oben beschrieben, mittels eines Tiefpasses das Signal zurückgewonnen. Die PPM/PDM-Wandlung kann so erfolgen, daß einem Flip-Flop (s. Abschn. 5.4.) synchron die kurzen Abtastimpulse und die PPM-Impulsfolge zugeführt werden. Durch die Abtastimpulse erscheint am Flip-Flop-Ausgang der Einszustand, der durch die PPM-Impulse wieder den Nullzustand annimmt. Bild 6.50 verdeutlicht, daß dabei eine dauermodulierte Impulsfolge entsteht. Im anschließenden Tiefpaß wird die ursprüngliche Sendeschwingung wiedergewonnen.

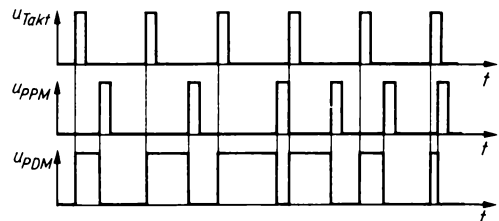


Bild 6.50. Prinzip der PPM-PDM-Wandlung

Zusammenfassung

Zur Amplitudenmodulation ist ein Bauelement mit nichtlinearer Strom-Spannungs-Kennlinie erforderlich. Für kleine Leistungen werden Dioden- und Transistorschaltungen, für wesentlich größere entsprechende Elektronenröhrenschaltungen angewendet. AM-Modulatoren werden entweder als Eintakt- oder als Gegentaktschaltung realisiert. Der Ringmodulator ist eine Doppelgegentschaltung mit großer praktischer Bedeutung, weil bei ihr bei symmetrischem Aufbau alle Schwingkreise gegeneinander entkoppelt sind und nur wenig unerwünschte Seitenfrequenzschwingungen höherer Ordnung auftreten.

Die Signalerückgewinnung heißt Demodulation. Diese wird bei relativ niedriger Trägerfrequenz und kleiner Bandbreite durch Diodenschaltun-

gen vorgenommen. Grundsätzlich können dabei Diode und Belastungswiderstand in Reihe oder parallel geschaltet werden. Durch einen Kondensator wird die Diode mit der Amplitude des modulierten Trägers vorgespannt, und es erfolgt eine Spitzengleichrichtung.

Bei höheren Trägerfrequenzen und größerer Bandbreite (z. B. zwischenfrequente AM-Bildsignale beim Fernsehen) wird gegenwärtig ein Synchron- oder Produktdemodulator eingesetzt. In diesem wird die zu demodulierende AM-Schwingung durch die in einem Begrenzerverstärker gewonnene Trägerfrequenzimpulsfolge durchgeschaltet. Das führt im Ausgangskreis zu einer Stromänderung, die der Modulationsspannung proportional ist.

Bei der Winkelmodulation wird der Winkel der Trägerschwingung im Rhythmus der Sendefunktion verändert; die Amplitude bleibt konstant. Man unterscheidet zwischen Frequenz- und Phasenmodulation.

Bei einem Frequenzmodulator wird die Eigenfrequenz des selbstschwingenden Oszillators durch einen von der Modulationsspannung steuerbaren Blindwiderstand verändert.

Bei einem Phasenmodulator bleibt die Oszillatorfrequenz konstant; durch einen von der Modulationsspannung steuerbaren Blindwiderstand wird lediglich ein im Ausgang liegender Schwingkreis verstimmt und damit der Phasenwinkel der Oszillatorschwingung verändert. Als steuerbare Blindwiderstände werden neben Kapazitätsdiodenschaltungen vor allem Reaktanzschaltungen eingesetzt.

Zur Demodulation frequenz- und phasenmodulierter Schwingungen werden die gleichen Verfahren angewendet. Die entsprechenden Schaltungen werden in der Praxis als FM-Demodulatoren bezeichnet.

Es gibt vier unterschiedliche Verfahren, um winkelmodulierte Schwingungen zu demodulieren:

- FM-AM-Wandlung, Signalrückgewinnung durch AM-Demodulator
- FM-PDM-Wandlung, Signalrückgewinnung durch Tiefpaß
- FM-PFM-Wandlung, Signalrückgewinnung durch Tiefpaß
- Demodulation durch Phasenregelung.

Die nach diesen Verfahren arbeitenden Schaltungen sind amplitudenabhängig. Es muß also stets durch vorgeschaltete Begrenzerstufen für eine konstante Amplitude gesorgt werden.

Die Pulsmodulation und -demodulation kann mit vielen unterschiedlichen Schaltungen durchgeführt werden. Voraussetzung für alle Verfahren ist eine Pulsamplitudenmodulation. Für die PAM ist ein elektronischer Schalter erforderlich, der durch einen Impulsgenerator gesteuert wird. Die angelegte Sendespannung wird dann abgetastet und in eine amplitudenmodulierte Impulsfolge umgewandelt. Durch nachgeschaltete Digitalschaltungen kann die PAM- in eine PDM-, PPM- oder PFM-Impulsfolge übergeführt werden. Größeren Aufwand erfordert die PCM, da nach der notwendigen PAM eine Quantisierung und anschließende Kodierung der Impulse erfolgen muß.

Die Demodulation kann bei PAM, PDM und PFM mit einem Tiefpaß vorgenommen werden. Bei PPM und PCM muß vorher eine Wandlung in PAM oder PDM erfolgen.

6.3. Aufgaben

1. Erläutern Sie, warum bei der Nachrichtenübertragung die Signalschwingung einem Träger aufgeprägt wird!
2. Vergleichen Sie die Modulationsverfahren AM, FM und PM bezüglich der
 - Beeinflussung der in der Grundgleichung enthaltenen Größen einer Trägerschwingung
 - Modulationsintensität
 - Bandbreite des Frequenzspektrums
 - Leistungsbilanz!
3. Erläutern Sie, warum bei einer Frequenzmodulation ($\Delta\Omega = \text{konst.}$) der Phasenhub umgekehrt proportional der Signalfrequenz ist! Stellen Sie dazu zwei Liniendiagramme einer frequenzmodulierten Trägerschwingung dar, und zwar mit einer tiefen und mit einer hohen Signalfrequenz!
4. Erläutern Sie, warum bei der Übertragung breitbandiger Signale durch Winkelmodulation des Trägers die hohen Modulationsfrequenzen senderseitig akzentuiert (Preemphasis) und empfängerseitig deakzentuiert (Deemphasis) werden!
5. Erläutern Sie das Prinzip der zeitmultiplexen Nachrichtenübertragung bei den verschiedenen Pulsmodulationsarten!
6. Begründen Sie, warum jede AM-Schaltung ein Modulationsbauelement mit nichtlinearer Strom-Spannungs-Kennlinie erfordert!
7. Erläutern Sie das Verfahren der Diodendemodulation, und stellen Sie die Eigenschaften der Reihen- und Parallelschaltung gemäß Bild 6.19 gegenüber!

8. Erklären Sie die Wirkungsweise der Schaltstufe des Synchrondemodulators (Doppelgegentaktmischer)!
9. Erläutern Sie die unterschiedliche Wirkungsweise des im Bild 6.26 angegebenen Filterphasenmodulators gegenüber dem im Bild 6.23 gezeigten Frequenzmodulator!
10. Nennen Sie die grundsätzlichen Verfahren zur Demodulation frequenz- und phasenmodulierter Schwingungen!
11. Erklären Sie die Wirkungsweise eines Koinzidenzdemodulators mit gekoppelten Differenzverstärkern!
12. Erläutern Sie die prinzipiellen Verfahren zur Pulsmodulation!

Sachwörterverzeichnis

- A-Betrieb 80
- AB-Betrieb 80
- ADU 154
- Amplitudenbegrenzung 99
- Amplitudengang 59
- Amplitudenhub 164
- Amplitudenmodulation 163 f., 173
- Analog-Digital-Signalwandlung 93
- Analog-Digital-Umsetzer 154
- Anstiegsverzögerung 108
- Anstiegszeit 64
- Arbeitspunkt 72
- Arbeitspunktstabilität 72
- Arbeitswiderstand 38
- Astabiler Kippgenerator 109
- Asynchrone Zähler 146
- Ausgangslastfaktor 120
- Ausgangsmodulation 176
- Ausgangsoffsetspannung 61
- Ausgangspuffer 127
- Ausgangsspannungshub 55
- Ausgangsstrombegrenzung 88
- Ausgangswiderstand 38
- Balance 97
- Bandbreite 165, 168 f.
- Bandfilter 76
- bandfiltergekoppelter Verstärker 76
- Bandgap 46
- Bandgap-Referenzquelle 22
- Basisbreitensteuerung 93
- Basisschaltung 39
- Basisvorspannung 72
- B-Betrieb 80
- BCD-zu-7-Segment-Dekoder 151
- Begrenzverstärker 178
- Bessel-Funktion 167
- Bessel-Funktionswerte 169
- Betriebsgröße 50
- Betriebsspannungsunterdrückung 64
- Biasstrom 63
- BIFET-Operationsverstärker 68
- Binärkode 171
- Bipolartransistor 43
- Bode-Diagramm 60
- Bootstrap-Schaltung 40, 111
- Bottstap-Schaltung 86
- Boucherot-Glied 91
- Breitbandfrequenzmodulation 170
- Breitbandverstärker 75
- Brückenschaltung 7
- Brückenverstärker 92
- Brummodulation 38
- Brummspannungsunterdrückung 23
- Bus 122
- C-Betrieb 80
- Clamping-Diode 123
- CMOS 126
- CMOS-HCT 129
- Colpitts-Oszillator 101
- commonmode-rejection 63
- Darlington-Schaltung 41
- Darlingtonschaltung 145
- DAU 154, 159
- D-Verstärker 93
- Deemphase 170
- Demodulation 163, 177
- Dezimalzähler 147
- Differenz-Pulskodemodulation 172
- Differenzfrequenz 164
- Differenzierglied 109
- Differenzverstärker 21
- Differenzverstärkerschaltung 43
- Differenzverstärkung 44
- Digit 155
- Digital-Analog-Umsetzer 154
- Diodendemodulator 177
- Diodenschaltung 175
- Direktgekoppelter Verstärker 77
- Diskriminator 182
- Doppelgegentaktmodulation 176
- Doppelgegentaktschaltung 75
- Doppel-T-Glied 78
- Dreipunktoszillator 101
- dual slope integrating AD-converter 157
- Dualzähler 147
- dynamischer Kennwert 64
- ECL-Gatter 129
- Eingangsbiasstrom 63
- Eingangsfehlerspannung 63
- Eingangsmodulation 176
- Eingangsoffsetspannung 63
- Eingangsoffsetstrom 63
- Eingangswiderstand 38
- Einseitenband-Amplitudenmodulation 166
- Eintaktverstärker 81
- Einweggleichrichtung 7
- Elektrometerverstärker 57
- elektronische Sicherung 19
- elektronische Stabilisierungsschaltung 17
- elektronischer Schalter 27, 117
- elektronisches Einstellglied 96
- Emitterkombination 75
- Emitterschaltung 39
- Endverstärker 37
- Entladezeitkonstante 9
- Erholzeit 139
- Festspannungsregler 21
- Filterphasenmodulator 181
- Flip-Flop 133
- Floatingregler 21
- FM-AM-Wandlung 181
- FM-DEM-Modulator 181
- FM-PDM-Wandlung 181
- FM-PFM-Wandlung 182
- Foldback-Kennlinie 20
- Frequenzaufbereitung 149
- Frequenzbandbreite 38

- Frequenzgang 56
 Frequenzgangkompensation 21
 Frequenzhub 166
 Frequenzkonstanz 100
 Frequenzmodulation 163, 166
 Frequenzmodulator 180
 Frequenzsynthese 117
 Frequenzteiler 149
 Frequenzumsetzung 163
 Frequenzverlauf 75
 FS 155
 full scale 155
 Funkentstörung 32

 Gegenkopplung 49
 Gegenkopplungsschaltung 51
 Gegentaktausgangsstufe 120
 Gegentaktbetrieb 47
 Gegentaktfankendemulator 182
 Gegentaktschaltung 7
 Gegentakstverstärker 82
 gehörrihtige Lautstärkeneinstellung 95
 Generator 99
 Gesamtleistung 165, 169
 Gleichrichterschaltung 7
 Gleichstromgegenkopplung 14
 Gleichakteingangswiderstand 58
 Gleichtaktspannung 58
 Gleichtaktunterdrückung 58
 Gleichtaktunterdrückung CMR 44
 Gleichtaktverhalten 56
 Gleichtaktverstärkung 44

 Harmonische 81
 Hartley-Oszillator 102
 H-Ersatzschaltung 50
 Hilbert-Transformation 60
 Höheneinsteller 95
 H-Parameter 50

 I2L-Gatter 130
 ideale OPV 55
 Impedanzwandler 57
 Impulsdiagramm 133
 Impulsfolge 93
 Impulsverzögerung 144
 integrierter Schaltkreis 89
 Interface-Schaltung 145
 Interienglied 108, 111
 invertierender Eingang 21
 invertierender Verstärker 56

 Kaskadenschaltung 41
 Kaskodeschaltung 42
 Kehrlege 165
 Klangbild 97
 Klangeinsteller 95
 Klirrfaktor 38

 Koinzidenzdemodulator 185
 Koinzidenzschaltung 184
 Kollektorschaltung 40
 Kombinationston 81
 kombinatorische Schaltung 118
 Komparator 156ff.
 Komplementär-Darlington-Schaltung 42
 Konstantstromausgang 152
 Konstantstromquelle 45, 111, 158
 Konstantstromsenken 152
 Koppelschaltung 47
 Kopplungsvierpol 73
 kreuzgekoppelter Differenzverstärker 179

 Ladekondensator 10
 Langeinstellnetzwerk 95
 last significant digit 155
 Lasttransistor 125
 latch-up-Effekt 64, 128
 Lautsprecherkurzschlußschutz 91
 Lautstärke 97
 LC-Siebglied 11
 least significant bit 155
 Leerlaufspannungsverstärkung 63
 Leistungsbilanz 166, 170
 Leistungsoperationsverstärker 70
 Linearitätsfehler 155
 LSB 155
 LSD 155

 Mehrflankenintegrationsverfahren 158
 mehrgliedrige Siebkette 13
 Meißner-Oszillator 100
 Miller-Integrator 112, 158
 mitlaufende Ladespannung 111f.
 Mittenspannungsteiler 91
 Modulation 163
 Modulationsgrad 164
 Modulationsindex 166
 Modulationstrapez 165
 Modulator 93
 MOS-Gatter 124
 Most significant bit 155
 Most significant digit 155
 MSB 155
 MSD 155
 Multiplexer 152
 Multivibrator 109

 Negativregler 22
 Netzfilter 33
 Netzteil 91
 nichtinvertierender Verstärker 57
 nichtinvertierender Eingang 21
 Norton-Verstärker 70
 Nulldurchgangsdiskriminator 186
 Nullpunktfehler 61

 Oberwellen 81
 offene Spannungsverstärkung 55
 Offset 56
 Offsetfehler 155
 Offsetkompensation 61
 Offsetspannung 55
 Offsetstrom 63
 open collector output 121
 open-collector 66
 Operationsverstärker 53, 105
 optoelektronischer Koppler 79
 Oszillator 99

 Parallel-Gegentaktschaltung 82
 Parallel-Parallel-Gegenkopplung 51
 Parallel-Serien-Gegenkopplung 51
 Parallel-Serien-Wandlung 152
 Parallele Eingabe 159
 Paralleleingabe 151
 Parallelregelung 18
 Parallelspeisung 101
 Parallelwandler 159
 Pegel 116
 Phasendemodulator 184
 Phasendiskriminator 182
 Phasengang 59
 Phasenhub 168, 181
 Phasenmodulation 163, 168
 Phasenmodulator 181
 Phasenrand 60
 Phasenregelung 182, 186
 Phasenschieber 184
 Phasensicherheit 60
 Phasenteilheit 184
 Phasenumkehrstufe 84
 Phasenvergleichstufe 186
 Phasenverschieberkette 103
 physiologische Lautstärkeeeinstellung 95
 Positivregler 22
 Präzisionsoperationsverstärker 69
 Preemphase 170
 Priorität 148
 Produktdemodulator 178
 Programmierbarer OPV 69
 pull down 121, 127
 pull up 121, 127
 pull-up-Widerstand 132
 pull-up-Widerstand 146
 Pulsamplitudenmodulation 170, 187
 Pulsdauermodulation 24, 171, 187
 Pulsfrequenzmodulation 24, 171, 188
 Pulsierende Gleichspannung 7
 Pulsodemodulation 171
 Pulsmodulation 163, 170
 Pulsphasenmodulation 171, 187

- Quadraturdemodulator 184
- Quantisierungsrauschen 171
- Quantisierungsstufe 171
- Quarzoszillator 102
- quasikomplementärer Gegentaktverstärker 88
- 2R/R-Netzwerke 160
- Ratiodetektor 183
- RC-Oszillator 103
- RC-Siebglied 11
- RC-Verstärker 73
- RC-Vorverstärker 76
- Reaktanzschaltung 180f.
- Reale OPV 55
- Rechenverstärker 53
- Referenzspannung 158
- Referenzelement 46
- Referenzspannung 18, 158
- Referenzspannungsquelle 45
- Regellage 165
- Reihenregelung 17
- Richtfunktechnik 172
- Riegger-Kreis 182
- Ringmodulator 175
- RS-Flip-Flop 133
- Rückkopplung 49
- Rückkopplungsgrad 51
- Rückkopplungsgrundgleichung 51
- Rückwandlung 93
- Ruhestromstabilisierung 85
- Sägezahngenerator 111
- Sägezahnumsetzer 156
- Sättigungsbegrenzung 28
- Schaltnetzteil 31
- Schaltregler 17
- Schaltstufe 178, 184
- Schleifenverstärkung 51
- Schmalbandverstärker 76
- Schutzschaltung 54
- Schwingkreis 100
- 7-Segment-Anzeige 151
- Segmenttreiber 153
- Seitenfrequenzschwingung 164
- Selbsterregungsbedingung 99
- Selbstinduktionsspannungsspitze 118
- Selektivverstärker 76
- Sendefunktion 163
- Senderverstärker 37
- sequentielle Schaltung 133
- serielle Eingabe 151, 159
- Serien-Gegentaktsschaltung 82
- Serien-Parallel-Gegenkopplung 51
- Serien-Serien-Gegenkopplung 51
- Serienspeisung 101
- Setstrom 69
- Siebfaktor 12
- Siebschaltung 11
- Signal-Rausch-Abstand 170
- Signallaufzeit 144
- Sinusgenerator 100
- slew-rate 64
- soft-recovery-Verhalten 25
- soft-recovery-Charakteristik 9
- SOAR-Schutz 91
- Spannungs-Frequenz-Umsetzung 159
- Spannungs-Zeit-Umsetzung 156
- Spannungsfolger 57
- Spannungsgegenkopplung 51
- spannungsgesteuerter Oszillator 186, 188
- Spannungspegel 116
- Spannungsstabilisierung 15
- Spannungsverdopplerschaltung 9
- Spannungsverstärker in offener Schleife 55
- Spannungsverstärkung 63
- Spannungsverstärkung in offener Schleife 63
- Spektraldiagramm 165, 167, 169
- Sperrwandler 27
- Spitzengleichrichtung 177
- Stabilisierungsfaktor 15
- Stabilisierungsschaltung mit Z-Diode 15
- Stand-by 93
- Standard-Darlington-Schaltung 42, 85
- statische Störsicherheit 123
- statistischer Kennwert 63
- Stellentreiber 153
- Stelltransistor 17
- Stereo-Endverstärker 93
- Stereobalanceeinsteller 96
- Stereobasisbreite 93
- stetiger Regler 17, 24
- Störabstand 38
- Störkapazität 124
- Stromdifferenzeingang 70
- Stromflußwandler 27
- Stromflußwinkel 10
- Stromflußwinkelmodulation 175
- Stromgegenkopplung 51
- Stromquellenbank 45
- Stromspiegel 45
- Stromstabilisierung 15
- Stromsteuerprinzip 129
- Stromsteuerung 101
- sukzessive Approximation 158
- Summenfrequenz 164
- Summierwandler 27
- supplyvoltage-rejektion 64
- Synchrodemodulator 178
- synchrone Zähler 146
- systemfremde Last 145
- Tastverhältnis 137
- Temperaturschutz 91
- three state output 122, 128
- Tiefeneinsteller 95
- totem pole output 121
- Trägerleistung 165
- Trägerschwingung 163
- Trägerstauereffekt 9
- Transferkennlinie 123, 128
- Transistorausgangsmodulator 176
- Transistoreingangsmodulator 176
- Transitfrequenz 55
- Transverter 26
- Treiberstufe 84
- Treibertransistor 125
- Trigger 133
- TTL-Gatter 119
- Überschwingen 38
- Überspannungsschutz 91
- Übersteuerungstechnik 120
- Übertragsimpuls 148
- Übertragungsverhalten 74
- Umsetzrate 156
- Umsetzzeit 155
- VCO 102, 159
- VCXO 103
- Verhältnisdiskriminator 183
- Verlustleistungshyperbel 82
- Verstärker 37
- Verstärkungseinsteller 95
- Verstärkungsfehler 155
- Verstärkungsgrad 38
- Verweilzeit 139
- Verzerrung 38
- Vierpol 49
- Vierquadrantendemodulator 184
- virtuelle Masse 55
- voltage controlled oscillator 102
- Vorverstärker 37
- Wahrheitstafel 118
- Wechselrichtung 27
- Widlar-Diode 47
- Wien-Brücke 103
- Wien-Brücken-Oszillator 104
- Winkelmodulation 163
- Zähldiskriminator 186
- Zeigerdiagramm 165, 167, 169
- Zeitkonstante 24, 107
- Zeitmultiplexverfahren 170, 172
- Zweiflankenintegrationsverfahren 157
- Zweiseitenband-Amplitudenmodulation 166
- Zweiweggleichrichtung 7

ISBN 3-341-00564-1