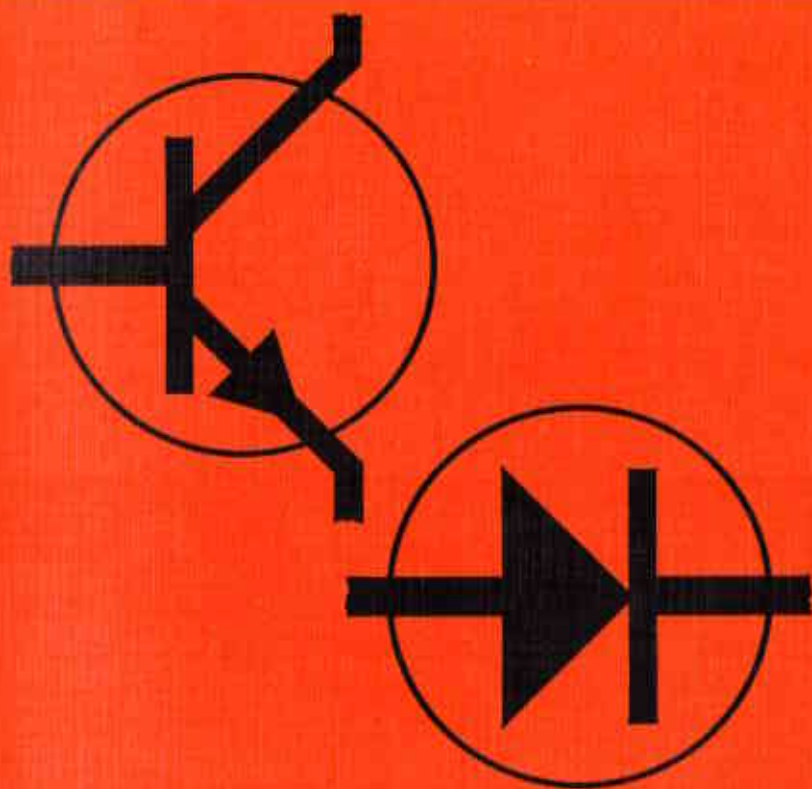


Grundlagen der Elektronik



Repetitor „Grundlagen der Elektronik“

Das erworbene Wissen kann immer nur durch Wiederholungen, erneutes Durcharbeiten und laufende Erfolgskontrollen vertieft und gefestigt werden. Bei dieser wichtigen Aufgabe leistet der Repetitor „Grundlagen der Elektronik“ wertvolle Dienste. In diesem Band wird der gesamte Lehrstoff des vorgenannten Handbuchs — systematisch programmiert — anhand einer Vielzahl von stoffbezogenen Fragen wiederholt. Jede Frage wird ausführlich beantwortet; hierbei werden allgemeingültige Begriffe als Merksätze einprägsam herausgestellt.

Der Aufbau des Repetitors, der in methodischer und didaktischer Hinsicht neue Maßstäbe setzt, erleichtert die Lernarbeit ganz wesentlich. Das Handbuch und der Repetitor sind eng miteinander verzahnt und stellen ein Ganzes dar. Beide Lehrwerke ermöglichen es dem Leser, sich in kurzer Zeit mit den Grundlagen der Elektronik eingehend vertraut zu machen.

Handbuch der Elektronik

Teil 1 — Analogtechnik

In diesem Lehr- und Lernbuch wird folgender Lehrstoff einprägsam und anschaulich behandelt:

Kurze Wiederholung der elektrotechnischen Grundlagen
Physikalische Grundlagen der Halbleiter
Halbleiterdioden und ihre Anwendung
Transistorgrundschaltungen
Der Transistor als Verstärker
Der Transistor als Schwingungserzeuger
Vierschicht-Halbleiter
Unijunctionstransistor
Stromversorgungsschaltungen mit Transistoren
Stromversorgungsschaltungen mit Thyristoren und TRIACs
Fotoelektronische Bauelemente
Feldeffekt-Transistoren

Zahlreiche Abbildungen, Kennlinien und Rechenbeispiele ergänzen den Lehrstoff.

Teil 2 — Digitaltechnik

Anschließend an das Stoffgebiet der Analogtechnik wird im Teil 2 des „Handbuchs der Elektronik“ die Digitaltechnik mit folgenden Schwerpunkten behandelt:

Grundlagen der Digitaltechnik
Verknüpfungsglieder
Impulsformer
Kippschaltungen
Schaltwerke
Codewandler
Datenübertragungstechnik
Magnetkern-technik
Grundsätzliches über EDV-Anlagen
Aufbau elektronischer Schaltkreise

Zahlreiche Abbildungen, Schaltungen und Rechenbeispiele ergänzen den Lehrstoff.

Grundlagen der Elektronik

5., verbesserte Auflage

Herausgeber: Institut zur Entwicklung moderner Unterrichtsmedien e. V.
28 Bremen 1 — Bahnhofstraße 10

Vorwort

Der Begriff Elektronik dürfte für den Leser zunächst noch mit vielen Fragezeichen versehen sein. Dabei bedeutet Elektronik weiter nichts als das Regeln und Steuern technischer Vorgänge mit Hilfe elektrischer Mittel, wobei heute Schaltungen mit Halbleiter-Bauelementen im Vordergrund stehen. Mit dem vorliegenden Band möchten wir die Grundlagenkenntnisse vermitteln, die zur weiteren Einarbeitung in das Gebiet Elektronik unerlässlich sind. Dabei setzen wir voraus, daß die Grundlagen der allgemeinen Elektrotechnik, also die Zusammenhänge im Gleich- und Wechselstromkreis, bekannt sind.

Dieser Band ist in erster Linie als Lehr- und Lernbuch gedacht; er enthält aber auch eine Reihe von Schaltungs- und Bauvorschlägen, die den Leser zu eigener, praktischer Beschäftigung mit der Elektronik anregen sollen. Da im Mittelpunkt eines Unterrichtsgesprächs eigentlich immer die Frage nach dem Warum steht, haben wir auf die Darstellung physikalischer Zusammenhänge nicht ganz verzichtet. Die Erläuterungen sind jedoch einfach und verständlich gehalten.

Das vorliegende Lehrbuch wird durch den Repetitor „Grundlagen der Elektronik“ ergänzt, der Gelegenheit bietet, den erarbeiteten Lehrstoff — systematisch programmiert und vielfach verzweigt — in Frage und Antwort zu wiederholen. Wir hoffen, mit dem vorliegenden, aus der Praxis für die Praxis geschriebenen Band eine Lücke in der Ausbildung zum Elektroniker zu schließen.

Die Herausgeber

Stand: Frühjahr 1974

Nachdruck, auch auszugsweise, nicht gestattet.

Inhaltsverzeichnis

	Seite
1. Meßtechnik	
1.1. Welche Bedingungen müssen Zeigerinstrumente für Messungen in der Elektronik erfüllen?	11
1.2. Messungen mit Zeigerinstrumenten	13
1.2.1. Spannungsmessungen mit Zeigerinstrumenten	14
1.2.2. Strommessungen mit Zeigerinstrumenten	15
1.2.3. Widerstandsmessungen mit dem Ohmmeter	17
1.2.4. Gleichzeitige Spannungs- und Strommessung	18
1.2.5. Indirekte Strommessung	19
1.2.6. Die Messung von Wechselspannungen und -strömen mit Zeigerinstrumenten	20
1.3. Das Messen mit Oszilloskopen (Oszillografen)	21
1.3.1. Grundsätzlicher Aufbau und Wirkungsweise eines Oszilloskops	21
1.3.2. Der elektronische Schalter als Zusatzgerät	32
1.3.3. Einfache Messungen mit dem Oszilloskop	33
2. Halbleiter	
2.1. Was ist ein Halbleiter?	42
2.1.1. Kristallaufbau der Halbleiterstoffe	42
2.1.2. N-Leiter/P-Leiter (Störstellenleiter)	44
2.1.3. Das Temperaturverhalten von Halbleitern	46
2.1.4. Die Berücksichtigung der zulässigen Verlustleistung eines Bauelements im Kennlinienfeld	50
2.2. Einfache Halbleiter-Bauelemente	53
2.2.1. NTC-Widerstände (Heißleiter/Thermistoren)	53
2.2.2. PTC-Widerstände (Kaltleiter)	57
2.2.3. Hall- oder Feldplatten (Hallsonden)	60
2.2.4. Fotowiderstände	63
3. Halbleiterdioden (Zweischicht-Halbleiter)	
3.1. Grundsätzlicher Aufbau und Wirkungsweise einer Diode	66
3.1.1. Plus am P-Leiter; Minus am N-Leiter	66
3.1.2. Minus am P-Leiter; Plus am N-Leiter	67
3.1.3. Die Diffusionsspannung	68
3.1.4. Die Kennlinie einer Diode (PN-Übergang)	70
3.1.5. Statischer und dynamischer Widerstand	71
3.2. Messungen an Dioden	74
3.2.1. Einfache Messungen mit dem Ohmmeter	75
3.2.2. Die Kennlinien verschiedener Dioden	76

	Seite
3.3. Universaldioden und ihre Anwendung	78
3.3.1. Die Diode als Gleichrichter	78
3.3.2. Die Diode als Schalter	84
3.3.3. Die Diode als Entkoppler	87
3.3.4. Die Diode als Spannungsbegrenzer	89
3.3.5. Die Diode als Leistungsbegrenzer	93
3.3.6. Die Diode als Gegenzelle	95
3.4. Referenzdioden (Z-Dioden)	96
3.4.1. Grundsätzlicher Aufbau und Wirkungsweise	96
3.4.2. Spannungs-Stabilisierungsschaltungen mit Referenzdioden	99
3.4.3. Weitere Anwendungsmöglichkeiten für Referenzdioden	107
3.5. Varistoren (VDR-Widerstände)	109
3.5.1. Grundsätzlicher Aufbau und Wirkungsweise	109
3.5.2. Anwendung von Varistoren	111
3.6. Fotodioden	112
3.6.1. Grundsätzlicher Aufbau und Wirkungsweise	112
3.6.2. Anwendung von Fotodioden	114
3.7. Lumineszenzdioden (LED)	115
3.7.1. Grundsätzlicher Aufbau und Wirkungsweise	115
3.7.2. Anwendung von Lumineszenzdioden	118
3.8. Kapazitätsdioden (Varaktoren)	119

4. Der Transistor

4.1. Grundsätzlicher Aufbau und Wirkungsweise	122
4.1.1. Bauformen und Anschlußarten der Transistoren	128
4.1.2. Die Bezeichnung und Zählrichtung der Ströme und Spannungen eines Transistors	131
4.2. Einfache Messungen an Transistoren	132
4.2.1. Ermittlung der Schichtfolge (PNP oder NPN) eines Transistors und Prüfung seiner Funktionsfähigkeit	132
4.2.2. Messung des Gleichstromverstärkungsfaktors B	136
4.3. Kennlinien des Transistors (Emitterschaltung)	138
4.3.1. Die Eingangskennlinie eines Transistors	138
4.3.2. Die Ausgangskennlinien eines Transistors	141
4.3.3. Die Steuerkennlinien eines Transistors	144
4.4. Der Transistor als Verstärker	148
4.4.1. Die wichtigsten Eigenschaften eines Transistors als Verstärker	150
4.4.2. Die drei Grundschaltungen für einen Transistor	156
4.4.3. Die Stromversorgung von Transistorverstärkern	159
4.4.4. Die Ein- und Auskopplung des Signals	162
4.4.5. Die Arbeitspunktstabilisierung der Transistorverstärker	167
4.4.6. Die richtige Bemessung des Lastwiderstands anhand des Ausgangskennlinienfeldes	174

	Seite
4.4.7. Die Berechnung von Transistorverstärkern	179
4.4.8. Der praktische Aufbau eines Vorverstärkers	189
4.4.9. Mehrstufige Transistorverstärker	192
4.5. Der Transistor in Spannungs-Stabilisierungsschaltungen	197
4.5.1. Parallel-Stabilisierung der Ausgangsspannung	197
4.5.2. Serien-Stabilisierung der Ausgangsspannung	198
4.6. Der Transistor als Schalter	201
4.6.1. Die EIN-Stellung als Schalter	201
4.6.2. Die AUS-Stellung als Schalter	203
4.6.3. Vergleich zwischen einem elektromechanischen Schalter und dem Transistor als Schalter	207
4.7. Der Transistor als Schwingungserzeuger	209
4.7.1. Die Rückkopplung	209
4.7.2. Mitkopplung und Gegenkopplung	211
4.7.3. LC-Generator	212
4.7.4. RC-Generator (Wien-Robinson-Generator / Wiengenerator)	213
4.7.5. RC-Phasenkettengenerator	215
4.8. Sonderformen des Transistors	216
4.8.1. Der Feldeffekt-Transistor (FET)	216
4.8.2. Der Unijunktion-Transistor (UJT)	218
4.8.3. Der Fototransistor	220

5. Vierschichtbleiter-Bauelemente

5.1. Die Vierschichtdiode	222
5.2. Der Thyristor	223
5.2.1. Grundsätzlicher Aufbau und Wirkungsweise	223
5.2.2. Einfache Funktionsprüfung von Thyristoren	225
5.2.3. Anwendung von Thyristoren	226
5.3. Der DIAC (Triggerdiode)	228
5.4. Der TRIAC	229
5.4.1. Grundsätzlicher Aufbau und Wirkungsweise	229
5.4.2. Anwendung des TRIACs	230

6. Elektronenröhren

6.1. Grundsätzlicher Aufbau und Wirkungsweise der Elektronenröhre	232
6.1.1. Zweipolröhre (Diode)	234
6.1.2. Dreipolröhre (Triode)	234
6.2. Elektronenröhre als Verstärker	237
6.2.1. Arbeitspunkteinstellung bei Röhrenverstärkern	238
6.2.2. Verwendung der Röhre mit Steuergitter	240

	Seite
7. RC-Glieder	
7.1. RC-Glieder als Impulsformer	241
7.1.1. Elektrische Vorgänge beim Auf- und Entladen einer Kapazität	241
7.1.2. RC-Glied als Differenzierglied	243
7.1.3. RC-Glied als Integrierglied	245
7.1.4. Einfache Kipperschaltung mit Integrlerglied	247
7.1.5. RC-Glied als Beschleuniger	248
7.2. RC-Glieder als Hoch- und Tiefpaß	249
7.2.1. Differenzierglied als Hochpaß	250
7.2.2. Integrierglied als Tiefpaß	252
8. Kippstufen	
8.1. Allgemeines	254
8.2. Signalzustände beim Transistor als Schalter	255
8.3. Die bistabile Kippstufe (Flipflop)	256
8.3.1. Grundschialtung	256
8.3.2. Grundschialtung mit Hilfsspannung	260
8.3.3. Bistabile Kippstufe mit statischen Eingängen	261
8.3.4. Bistabile Kippstufe mit dynamischen Eingängen	265
8.3.5. Symbolische Darstellung des FF	267
8.3.6. Anwendungsbeispiele für FF	269
8.4. Monostabile Kippstufe / Schmitt-Trigger	272
8.5. Astabile Kippstufe	276
8.6. Selbstbau einer Kippstufe	277
9. Verknüpfungsglieder	
9.1. UND-Glied (UND-Gatter)	282
9.2. ODER-Glied (ODER-Gatter)	284
9.3. NICHT-Glied (NICHT-Gatter)	286
10. Anhang	
10.1. Die Bezeichnung von Halbleiter-Bauelementen	288
10.2. Die Bezeichnung von Elektronenröhren	290
10.3. Farbcode für Widerstände und Kondensatoren	292
10.4. Normreihen für Widerstände und Kapazitäten	294
10.5. Anschlußbeispiele für Dioden und Gleichrichter	296
10.6. Anschlußbeispiele für Transistoren	297
Sachwortverzeichnis	299

1. Meßtechnik

Elektronische Bauelemente und Schaltungen können nur durch Messungen auf ihre Funktionsfähigkeit und Eigenschaften geprüft werden. Kenntnisse über den grundsätzlichen Aufbau und die Funktion von Zeigerinstrumenten wie Spannungs-, Strom- und Widerstandsmesser werden vorausgesetzt. Im Rahmen dieses Bandes wird daher im wesentlichen auf meßtechnische Fragen eingegangen, die die Elektronik betreffen.

In der Elektronik wird mit den unterschiedlichsten Stromarten gearbeitet; die wichtigsten sollen daher zunächst erläutert werden. Dazu sei erwähnt, daß unter Stromart allgemein der zeitliche Verlauf sowohl des Stroms als auch der Spannung zu verstehen ist.

Gleichstrom: Stromart, bei der die Augenblickswerte für Strom und Spannung zeitlich konstant sind (gleichbleibende Stromrichtung). Man spricht auch dann von Gleichstrom oder Gleichspannung, wenn kleine, aber unwesentliche Schwankungen überlagert sind oder wenn infolge von Belastungsschwankungen Änderungen der an sich konstanten Größe auftreten (Abb. 1.1).

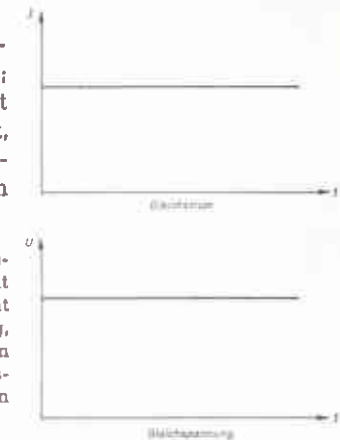


Abb. 1.1

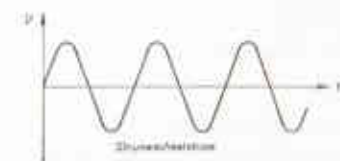
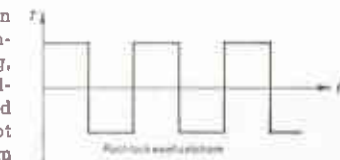


Abb. 1.2

Wechselstrom: Periodisch, d.h. in regelmäßigen Zeitabständen, im gleichen Verlauf wiederkehrender Strom (bzw. Spannung) wechselnder Richtung, aber beliebiger Kurvenform. Der lineare Mittelwert ist Null, d.h., die Summe aller positiven und negativen Augenblickswerte einer Periode ergibt den Wert Null. Sonderfall: Sinuswechselstrom bzw. Sinuswechselspannung.

Mischstrom: Entsteht durch die Überlagerung von Gleich- und Wechselstrom; deshalb ist der lineare Mittelwert nicht Null.

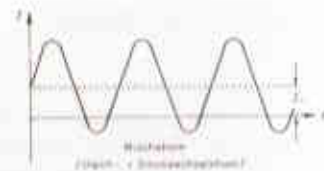


Abb. 1.3

Pulsstrom: Periodisch, d.h. in regelmäßigen Zeitabständen, wiederkehrender Strom- oder Spannungsstoß, wobei entweder immer die gleiche Richtung oder abwechselnd positive und negative Richtung auftritt. Bei dem durch Gleichrichtung aus Wechselstrom gewonnenen Pulsstrom spricht man in der Stromversorgungstechnik auch von „pulsierendem Gleichstrom“ bzw. von „pulsierender Gleichspannung“.

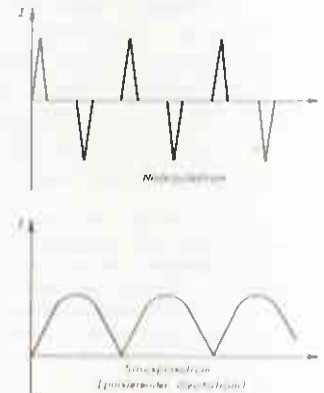


Abb. 1.4

Impuls: Kurzzeitig wirkender Strom- oder Spannungsstoß beliebiger Kurvenform; man unterscheidet zwischen einseitigem Impuls (ohne Richtungswechsel) und zweiseitigem Impuls (ein Richtungswechsel).

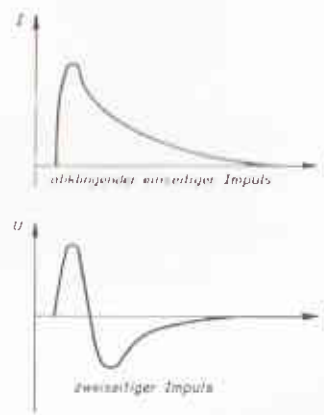


Abb. 1.5

Zeitabhängige Größen wie **Spannung** und **Strom** werden durch die nachstehend vermerkten **Formelzeichen** gekennzeichnet:

$U_- ; I_-$	Zeitlich konstante Gleichspannung oder Gleichstrom; wenn keine Verwechslung möglich, auch einfach U bzw. I .
$U_{eff} ; I_{eff}$	Effektivwert der Wechselspannung bzw. des Wechselstroms; wenn keine Verwechslung mit Gleichspannung oder Gleichstrom möglich, kann auch einfach U bzw. I gesetzt werden.
U_{max} oder $\hat{u}$	Höchstwert der Wechselspannung.
I_{max} oder $\hat{i}$	Höchstwert des Wechselstroms.
$u ; i$	Augenblickswert der Wechselspannung bzw. des Wechselstroms.
$U_{ss} ; I_{ss}$ (auch $\Delta U ; \Delta I$)	Spannungs- bzw. Stromunterschied zwischen dem höchsten und niedrigsten Augenblickswert (wird auch z.B. „ $U_{Spitze-Spitze}$ “ genannt).

1.1. Welche Bedingungen müssen Zeigerinstrumente für Messungen in der Elektronik erfüllen?

Für Messungen an Halbleiter-Bauelementen und elektronischen Schaltungen sind grundsätzlich nur Zeigerinstrumente mit eingebautem Drehspulmeßwerk geeignet. Dreheiseninstrumente scheiden wegen ihres zu hohen Eigenverbrauchs aus. Aber selbst äußerst empfindliche Instrumente mit Drehspulmeßwerk sind zur Lösung vieler meßtechnischer Aufgaben der Elektronik ungeeignet, weil ihr Innenwiderstand als Spannungsmesser zu klein und als Strommesser zu groß ist. Hier die wichtigsten Anforderungen, die an Zeigerinstrumente zu stellen sind.

- Hohe Empfindlichkeit:** Die Meßgrößen wie Spannung und Strom haben vielfach so kleine Werte (mV; μ A), daß nur Meßinstrumente mit sehr hoher Empfindlichkeit genaue Messungen zulassen. Widerstände — z.B. Sperrwiderstände von PN-Übergängen — erreichen Werte bis zu einigen M Ω , so daß auch für Widerstandsmessungen eine sehr hohe Empfindlichkeit gefordert werden muß.
- Geringer Eigenverbrauch:** Unter Eigenverbrauch versteht man bei Zeigerinstrumenten mit Drehspulmeßwerk den zum Zeiger-

vollausschlag erforderlichen Strom. Im allgemeinen wird für Spannungsmesser heute der Kehrwert dieses Stroms ($\frac{1}{I}$) in der Einheit Ohm pro Volt ($\frac{\Omega}{V}$) angegeben. Aus dieser Angabe läßt sich für Spannungsmesser leicht der Innenwiderstand für einen bestimmten Meßbereich errechnen. Als Richtwert eines für Messungen in der Elektronik geeigneten Zeigerinstrumentes muß mindestens 10 000 Ohm pro Volt oder größer gefordert werden; das entspricht für Strommesser einem Eigenverbrauch von höchstens 100 μA oder weniger.

- c) **Hohe Anzeigegenauigkeit (Güteklasse):** Zur Aufnahme von Meßwertreihen zur Kennliniendarstellung müssen vielfach geringste Änderungen der Meßgröße eindeutig ablesbar sein. Dazu muß eine sehr genaue Anzeige der Meßgröße gefordert werden. Ein Zeigerinstrument sollte daher mindestens der Güteklasse 2 angehören (Anzeigefehler maximal $\pm 2\%$ vom Skalenendwert).
- d) **Mehrere Meßbereiche (Vielfachinstrument):** Die Kennlinien von Halbleiter-Bauelementen verlaufen mehr oder weniger stark gekrümmt; das bedeutet, daß im Spannungs-Strom-Verhalten eines solchen Bauelements große Änderungen auftreten. Diese können meßtechnisch nur dann genau erfaßt werden, wenn jeweils mehrere Meßbereiche für Spannungs-, Strom- und Widerstandsmessungen zur Verfügung stehen. Insbesondere bei der Aufnahme von Meßwertreihen wäre das Austauschen von Zeigerinstrumenten mit jeweils nur einem Meßbereich sehr aufwendig, zeitraubend und wegen des Bedarfs mehrerer Instrumente teuer. Aus diesen Gründen kommen zur Lösung meßtechnischer Aufgaben in der Elektronik nur Vielfachinstrumente in Betracht.

Abgesehen von den hohen Anforderungen, die an Zeigerinstrumente für Messungen in der Elektronik gestellt werden müssen, kann mit ihnen eine ganze Reihe meßtechnischer Aufgaben ohnehin nicht gelöst werden. Dazu gehören z.B. die Messung nichtsinusförmiger Wechselspannungen; die Aufnahme von Eingangskennlinien für Kleinleistungs-Transistoren oder die Messung des Isolationswiderstands am Eingang eines Feldeffekt-Transistors. Hierfür sind Oszilloskope, Röhren- bzw. Transistorvoltmeter oder andere Meßgeräte, wie Meßbrücken, einzusetzen. Im Rahmen dieses Bandes soll versucht werden, Mittel und Wege aufzuzeigen, wie man mit verhältnismäßig geringem Aufwand zu zumindest für die Beurteilung eines Bauelements oder einer Schaltung ausreichenden Ergebnissen kommen kann.



(Werkbild Hartmann & Braun AG)



(Werkbild Metrawatt)

Abb. 1.6 — Vielfachinstrumente

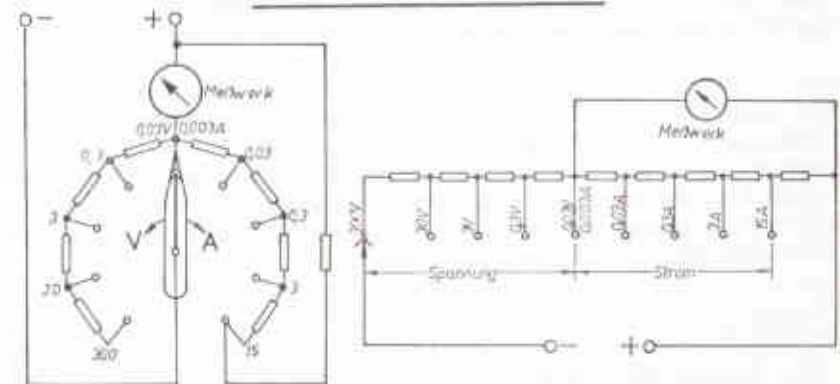


Abb. 1.7 — Schematischer Aufbau und Schaltung eines Vielfachinstruments für Spannungs- und Strommessungen

1.2. Messungen mit Zeigerinstrumenten

Zur richtigen Beurteilung der Meßergebnisse und Vermeidung grober Meßfehler ist die Kenntnis und Beachtung der Eigenschaften eines Meßinstrumentes oder einer Meßschaltung von großer Bedeutung.

1.2.1. Spannungsmessungen mit Zeigerinstrumenten

Ein Meßinstrument wird zur Spannungsmessung parallel zum Meßobjekt geschaltet. Das Instrument besitzt aber einen Innenwiderstand und verfälscht mit ihm das Meßergebnis. Soll z.B. die Teilspannung an einer Reihenschaltung von Widerständen nach nebenstehender Abbildung gemessen werden, so ändert sich das Verhältnis der Teilwiderstände durch das Anschalten des Spannungsmessers. Je kleiner der Innenwiderstand des Spannungsmessers ist, desto geringer wird der Spannungsabfall an der Parallelschaltung $R_2 \parallel R_{iVm}$. Daraus folgt: **Eine Spannungsmessung wird um so genauer, je größer der Innenwiderstand des Spannungsmessers ist.**

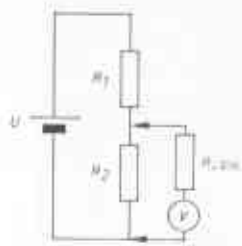


Abb. 1.8

Ideal wäre natürlich ein Spannungsmesser mit unendlich großem Innenwiderstand. Da es aber weniger auf den tatsächlichen Widerstand des Meßinstruments als vielmehr auf das Widerstandsverhältnis von R_{iVm} zu $R_{Meßobjekt}$ ankommt, kann eine einfache Faustregel angewandt werden: **Der durch das Anschalten eines Spannungsmessers an das Meßobjekt verursachte Fehler liegt bei 1 %, wenn der Innenwiderstand des Spannungsmessers mindestens 100mal so groß ist wie der Widerstand des Meßobjekts.**

Als Widerstand des Meßobjekts gilt dabei der Gesamtwiderstand, der zwischen den Anschlußklemmen des Spannungsmessers liegt. Je größer das Verhältnis von R_{iVm} zu $R_{Meßobjekt}$ gemacht werden kann, desto geringer wird der Schaltungsfehler. Bei nahezu allen modernen Vielfachinstrumenten läßt sich der Innenwiderstand für einen bestimmten Spannungsmeßbereich aus der Eigenverbrauchskonstante k in Ohm pro Volt ermitteln:

$$R_{iVm} = k \cdot U_{Vm}$$

Darin bedeuten:

R_{iVm} = Innenwiderstand des Spannungsmessers in Ohm,

k = Eigenverbrauch in Ohm pro Volt,

U_{Vm} = Spannungsmeßbereich (Skalenendwert) in Volt.

Beispiel:

Spannungsmesser mit einem Meßbereich 5 V, $k = 10\,000 \frac{\Omega}{V}$

$$R_{iVm} = 10\,000 \frac{\Omega}{V} \cdot 5 V = 50\,000 \Omega = 50 \text{ k}\Omega$$

Zur genauen Ablesung des Meßwerts ist wichtig, daß dieser möglichst am Ende des Skalenbereichs liegt. Der nach dem Klassenzeichen mögliche Fehler bezieht sich grundsätzlich auf den Skalenendwert und ist daher am Skalenanfang prozentual wesentlich größer.

Beispiel:

Spannungsmesser 100 V, Klasse 2 (d.h. Anzeigefehler $\pm 2 \%$ vom Skalenendwert).

Der Fehler kann demnach $\pm 2 \%$ von 100 V, also ± 2 V betragen. Bei einer Anzeige von nur 10 V im gleichen Meßbereich und dem Anzeigefehler von ± 2 V beträgt die Abweichung prozentual bereits $\pm 20 \%$!

Für die Wahl des richtigen Meßbereichs gilt also: **Meßbereich so wählen, daß die Anzeige des Meßwerts im letzten Drittel des Skalenbereichs liegt!**

Eines sollte dabei jedoch nicht außer acht gelassen werden: beim Umschalten des Spannungsmeßbereichs ändert sich der Innenwiderstand des Spannungsmessers! Das kann bei der Aufnahme von Meßwertreihen an einem bestimmten Meßobjekt zu unliebsamen Überraschungen führen. Da sich mit einem anderen Innenwiderstand auch ein anderes Verhältnis von R_{iVm} zu $R_{Meßobjekt}$ ergibt, tritt eine Änderung der Anzeigegenauigkeit auf. Für die Aufnahme von Meßwertreihen zur Kennliniendarstellung kann es daher manchmal von Vorteil sein, den Meßbereich nicht umzuschalten — selbst wenn die Anzeige dann nicht mehr im letzten Skalendrittel liegt. In solchen Fällen wählt man den Meßbereich nach dem höchsten zu erwartenden Meßwert und liest im gleichen Bereich auch kleine Meßwerte ab.

Beim Messen von Gleichspannungen ist auf die richtige Polarität zu achten. Die Anschlußklemmen oder -schnüre eines Spannungsmessers sind mit $\oplus$ und $\ominus$ gekennzeichnet. Sie sind so mit dem Meßobjekt zu verbinden, daß $\oplus$ am Pluspotential und $\ominus$ am Minuspotential liegt. Ist die Polarität der zu messenden Spannung unbekannt, so läßt sie sich mit Hilfe der Klemmenbezeichnungen ermitteln. Schlägt der Zeiger des angeschalteten Spannungsmessers zur richtigen Seite aus, so liegt die $\oplus$ -Klemme des Spannungsmessers am höheren Pluspotential der Meßschaltung. Dieses einfache Mittel kann man in vielen Schaltungen dazu benutzen, die Potentiale zu bestimmen. Bestehen z.B. Zweifel, wie ein gepolter Elektrolytkondensator in eine Schaltung einzufügen ist, so kann die entsprechende Polarität mit einem Spannungsmesser ermittelt werden.

1.2.2. Strommessungen mit Zeigerinstrumenten

Ein Strommesser ist in den Stromkreis — also in Reihe zum Meßobjekt — zu schalten. Der Innenwiderstand des Strommessers erhöht somit

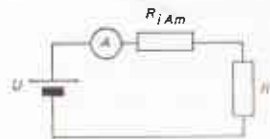


Abb. 1.9

den Stromkreiswiderstand. Durch das Einschalten eines Strommessers ergeben sich also auch Fehler. Nach dem Ohmschen Gesetz verringert sich in einem geschlossenen Stromkreis die Stromstärke, wenn der Stromkreiswiderstand vergrößert wird. Der Strommesser zeigt einen zu kleinen Wert an. Daraus folgt: **Eine Strommessung wird um so genauer, je kleiner der Innenwiderstand des Strommessers ist.** Für Strommesser wäre ein Innenwiderstand von 0 Ohm ideal. Da das nicht zu verwirklichen ist, muß man die Anforderungen hinsichtlich des Innenwiderstands der zu lösenden Meßaufgabe anpassen und kann dazu folgende Faustregel anwenden: **Der durch das Einschalten eines Strommessers in den Meßstromkreis bedingte Fehler beträgt etwa 1%, wenn der Innenwiderstand des Strommessers 1/100 des gesamten Stromkreiswiderstands beträgt.**

Die Bestimmung des Strommesser-Innenwiderstands ist schwieriger; in den meisten Fällen muß man ihn mit einem Widerstandsmeßinstrument messen.

Bei Vielfachinstrumenten ist manchmal folgendes Verfahren anwendbar. Besitzt das Vielfachinstrument einen Meßbereich, für den auf der Schalterskala eine Spannungs- und eine Stromangabe zu finden ist, so läßt sich aus diesen Angaben der Innenwiderstand für den empfindlichsten Meßbereich berechnen.

Beispiel:

Empfindlichster Meßbereich eines Vielfachinstruments 60 mV/0,6 mA.

Für diesen Bereich ergibt sich ein Innenwiderstand von

$$R_{iAm} = \frac{U_m}{I_m} = \frac{60 \text{ mV}}{0,6 \text{ mA}} = \underline{\underline{100 \Omega.}}$$

Zur Wahl des richtigen Meßbereichs gilt ebenso wie für Spannungsmesser: **Strommeßbereich so wählen, daß die Anzeige des Meßwerts im letzten Drittel des Skalenbereichs liegt!**

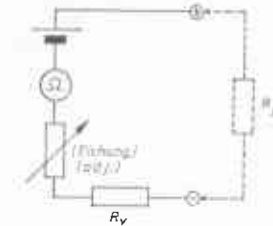
Bei der Aufnahme von Meßwertreihen kann es u.U. auch für Strommesser günstiger sein, den Meßbereich nicht umzuschalten; denn auch bei Strommessern ändert sich mit der Bereichsumschaltung der Innenwiderstand. Für Gleichstrommessungen ist die Beachtung der Polarität wichtig. Ein Strommesser wird als Verbraucher betrachtet und so in den Stromkreis geschaltet, daß sein mit $\oplus$ gekennzeichnete Anschluß am höheren Pluspotential der Schaltung liegt. Umgekehrt: Der Zeiger des Strommessers schlägt zur richtigen Seite hin aus, wenn über ihn der Gleichstrom von der $\oplus$ -Klemme zur $\ominus$ -Klemme fließt.

1.2.3. Widerstandsmessungen mit dem Ohmmeter

Ein Ohmmeter bietet die Möglichkeit, einfach und schnell die grundsätzlichen Eigenschaften und die Funktionsfähigkeit von vielen Halbleiter-Bauelementen zu prüfen oder zu beurteilen. Wir werden in den einzelnen Abschnitten wiederholt solche Prüf- und Meßbeispiele aufzeigen.

Das Ohmmeter besteht im Prinzip aus der Reihenschaltung einer Batterie (meistens Einzelzelle 1,5 V) mit einem in Ohm geeichten Strommesser. Gegebenenfalls können außerdem noch ein veränderbarer Widerstand zur Eichung und Vorwiderstände zur Einstellung verschiedener Meßbereiche in dieser Reihenschaltung liegen. Häufig benutzt man jedoch zur Eichung einen veränderbaren magnetischen Nebenschluß. Da eine Stromquelle eingebaut ist, stellt das Ohmmeter eine „aktive“ Prüfeinrichtung dar. Zur Bestimmung von Polaritäten ist wichtig:

Abb. 1.10 — Prinzipschaltung eines Ohmmeters



Für die Widerstandsmeßbereiche eines Vielfachinstruments liegt das Pluspotential der eingebauten Stromquelle an der mit $\ominus$ gekennzeichneten Anschlußklemme und sinngemäß das Minuspotential an der $\oplus$ -Klemme.

Bei einem einfachen Ohmmeter entsprechen die Klemmenbezeichnungen $\oplus$ und $\ominus$ der Batteriepolarität (Abb. 1.10).

Für die Elektronik sind solche Ohmmeter am besten geeignet, die die Messung sowohl kleiner Widerstände, z.B. Durchlaßwiderstände von PN-Übergängen, als auch sehr großer Widerstände, z.B. Sperrwiderstände von PN-Übergängen, zulassen; die entsprechenden Ohmmeter-Meßbereiche sind $\times 1$ bzw. $\times 10 \text{ k}$. Bei Messungen mit dem Ohmmeter sollte man nie vergessen, daß die Spannung der eingebauten Stromquelle u.U. nur kurzzeitig konstant bleibt. Für Widerstandsmessungen mit dem Ohmmeter ist daher die Beachtung folgender Regeln von großer Bedeutung:

- Vor jeder Messung ist ein Ohmmeter an den Klemmen kurzzuschließen und der Zeigerausschlag auf 0 Ohm zu eichen.
- Widerstandsmessungen sind so kurzzeitig wie möglich vorzunehmen, da sich sonst während der Messung die Spannung der eingebauten Stromquelle erheblich ändern kann.
- Die Anzeige eines Ohmmeters ist in der Skalenmitte am genauesten; der Meßbereich ist entsprechend zu wählen.

Datenblätter über Ohmmeter enthalten sehr häufig Meßbereichsangaben, die sich auf den Skalenendwert beziehen. Diese Angaben sind an sich nichtssagend, da die Skalenteilung im jeweiligen Endbereich keine genaue Ablesung eines Widerstandswerts zuläßt. Brauchbar sind dagegen Angaben über den in der Skalenmitte angezeigten Widerstandswert. Zur Messung sehr großer Widerstände sind in Ohmmeter vielfach besondere Batterien mit höherer Spannung (ca. 15 bis 20 V) eingebaut. Es kann sein, daß diese Spannung höher ist als die für ein Halbleiter-Bauelement zulässige Höchstspannung.

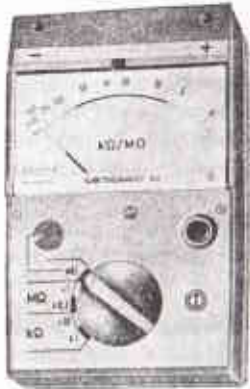


Abb. 1.11 — Ohmmeter

Vor Messungen im Hochohmbereich eines Ohmmeters sollte man sich daher davon überzeugen, ob die Spannung der eingebauten Batterie noch unterhalb der zulässigen Höchstgrenze für das zu messende Bauelement liegt. Ein Transistor würde z.B. durch eine zu hohe Ohmmeter-Spannung sofort zerstört; Abb. 1.11 zeigt ein modernes Ohmmeter.

Der Innenwiderstand eines Ohmmeters ist für den eingeschalteten Meßbereich gleich dem Widerstandswert in Skalenmitte.

Ein an sich selbstverständlicher Grundsatz soll in diesem Zusammenhang nicht unerwähnt bleiben: **Die Widerstandsmessung mit einem Ohmmeter ist nur möglich, wenn das Meßobjekt selbst strom- bzw. spannungslos ist.** Eine im Meßstromkreis zusätzlich vorhandene

Stromquelle verfälscht das Meßergebnis vollkommen oder führt gar zur Zerstörung des Meßinstruments.

1.2.4. Gleichzeitige Spannungs- und Strommessung

Die Lösung meßtechnischer Aufgaben zwingt in der Elektronik meistens zur gleichzeitigen Messung von Spannungen und Strömen. Solche Messungen können nur dann zu brauchbaren Ergebnissen führen, wenn bei ihrer Ausführung einige wichtige Zusammenhänge beachtet werden. Jetzt kommt es nämlich darauf an, das Verhältnis dreier Widerstände zueinander zu berücksichtigen: R_i des Spannungsmessers, R_j des Strommessers und R_x des Meßobjekts. Da der Widerstand von Halbleiter-Bauelementen vielfach stark von der angelegten Spannung abhängt, müssen u.U. während einer Meßreihe entsprechende Meßschaltungen geändert werden. Für die gleichzeitige Messung von Spannung und Strom an einem Meßobjekt gelten folgende grundsätzlichen Regeln:

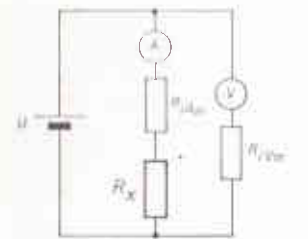


Abb. 1.12 — Schaltung zur Messung großer Widerstände

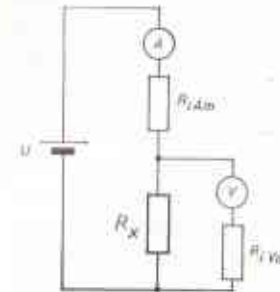


Abb. 1.13 — Schaltung zur Messung kleiner Widerstände

a) Ist der **Widerstand des Meßobjekts groß**, so wirkt sich der Widerstand eines Strommessers (Reihenschaltung) nur geringfügig auf den Gesamt Widerstand der Reihenschaltung aus; man wählt daher die nebenstehende Schaltung. In diesem Fall zeigt der Spannungsmesser zusätzlich den Spannungsabfall am Innenwiderstand des Strommessers an. Der Strommesser dagegen zeigt nahezu den richtigen Strom an, wenn R_x etwa 100mal so groß ist wie R_{iAm} .

b) Bei **kleinem Widerstand des Meßobjekts** ist die Schaltung zu ändern (Abb. 1.13). Der Strommesser zeigt hier zusätzlich den über den Spannungsmesser fließenden Strom an, dieser ist jedoch unter der Bedingung $R_{iVm} \geq 100 \cdot R_x$ so gering, daß das Meßergebnis ziemlich genau wird.

Vor dem Aufbau einer Meßschaltung ist unter allen Umständen zu prüfen, welche der beiden Meßschaltungen zweckmäßiger ist. Schon hier sei darauf hingewiesen, daß beide Bedingungen z.B. bei der Aufnahme einer Diodenkennlinie auftreten.

Im Durchlaßbereich wird eine Diode beim Überschreiten der Diffusionsspannung niederohmig. Dagegen ergeben sich im Sperrbereich im allgemeinen sehr hochohmige Widerstände. Eine durch eine Meßreihe (Spannung/Strom) aufgenommene Diodenkennlinie ist u.U. nur dann brauchbar, wenn man unter Berücksichtigung der Widerstandsverhältnisse beide Meßschaltungen nacheinander anwendet. (Näheres hierzu im Abschnitt 3 „Halbleiterdioden“.)

1.2.5. Indirekte Strommessung

Zur Strommessung in Schaltungen der Elektronik ist das Verfahren der indirekten Messung oft von großem Vorteil und sehr häufig sogar unumgänglich. Hierbei nutzt man die Tatsache aus, daß der Spannungsabfall an einem bestimmten ohmschen Widerstand unmittelbar von der Stärke des fließenden Stroms abhängt (Ohmsches Gesetz). **Indirekte Strommessung bedeutet die Messung des Spannungsabfalls an einem ohmschen Widerstand bekannter Größe.**

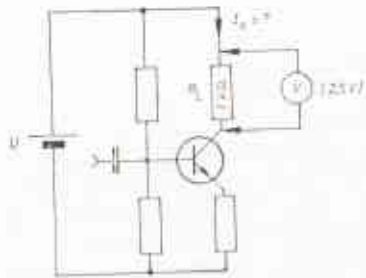


Abb. 1.14 — Indirekte Messung
des Kollektorstroms
einer Transistor-
schaltung

Beispiel: Anhand des nebenstehenden Schaltungsausschnitts ist der über einen Transistor fließende Strom I_C zu messen. Dazu müßte der Stromkreis aufgetrennt und ein Strommesser in Reihe geschaltet werden. Die Auftrennung des Stromkreises ist aber häufig sehr umständlich und vielfach sogar unmöglich; man braucht dabei nur an gedruckte oder integrierte Schaltungen zu denken. Ist jedoch der Widerstand R_L im Transistorstromkreis bekannt, so genügt die Messung des Spannungsabfalls an diesem Widerstand.

Aus dem gemessenen Spannungsabfall und dem bekannten Widerstandswert läßt sich der Strom nach dem Ohmschen Gesetz berechnen.

Beträgt der Spannungsabfall an R_L im dargestellten Schaltungsausschnitt 2,5 V, so muß der Strom I_C eine Stärke von

$$I_C = \frac{U_L}{R_L} = \frac{2,5 \text{ V}}{1 \cdot 10^3 \Omega} = 2,5 \cdot 10^{-3} \text{ A} = 2,5 \text{ mA} \text{ haben.}$$

Das Beispiel läßt erkennen, daß die indirekte Strommessung — genauer gesagt, die Berechnung des Stroms — um so einfacher wird, je runder der Widerstandswert ist. Wenn also der Widerstand frei wählbar ist, so sind Werte von 1 Ω , 10 Ω usw. am günstigsten.

Zur Kennliniendarstellung von Bauelementen an Oszilloskopen (Oszillografen) muß man zwangsläufig immer die indirekte Bestimmung der Stromstärke wählen, da Oszilloskope keine direkte Strommessung zulassen; wir kommen darauf noch zurück.

1.2.6. Die Messung von Wechselspannungen und -strömen mit Zeigerinstrumenten

Die Möglichkeit, Zeigerinstrumente zur Messung von Wechselspannungen oder -strömen in elektronischen Schaltungen einzusetzen, ist an sich auf wenige grundsätzliche Messungen beschränkt. Drehspulmeßwerke sind nur für Gleichstrommessungen geeignet; zur Wechselstrommessung werden Gleichrichter in Instrumente mit Drehspulmeßwerk eingebaut. Diese Gleichrichter setzen die Empfindlichkeit herab, was aus dem für Wechselstrom angegebenen Eigenverbrauch erkenntlich wird. So ergibt sich für ein und dasselbe Meßinstrument,

z.B. für Gleichspannungsmessungen 20 000 $\frac{\Omega}{V}$, dagegen für

Wechselspannungsmessungen 10 000 $\frac{\Omega}{V}$.

Nachteilig wirkt sich auch der Kennlinienknick des eingebauten Gleichrichters aus. Eine Gleichrichterwirkung tritt erst auf, wenn die anliegende Spannung größer als die Diffusionsspannung des eingebauten Gleichrichters ist. Daher sind Wechselspannungsmessungen nur für Werte oberhalb der Diffusionsspannung möglich. Die Skalenteilung ist jeweils am Anfang eines Wechselstrom- oder -spannungsbereichs enger.

Mit Zeigerinstrumenten lassen sich nur sinusförmige Wechselspannungen und Wechselströme messen, da ihre Skalen nach Effektivwerten geeicht sind. Ebenso ist der Frequenzbereich, innerhalb dessen genaue Meßergebnisse zu erwarten sind, verhältnismäßig eng begrenzt. Für Zeigerinstrumente mit eingebautem Drehspulmeßwerk und Gleichrichter kann ein **zulässiger Frequenzbereich von ca. 10 bis 10 000 Hz** angegeben werden. Das bedeutet, daß nicht einmal Wechselströme mit allen Frequenzen des Hörbereichs (16 bis 20 000 Hz) mit Zeigerinstrumenten meßbar sind. Hinzu kommt noch, daß Zeigerinstrumente keine meßtechnische Beurteilung hinsichtlich Verzerrungen von Wechselspannungen oder -strömen zulassen. Zusammenfassend kann also gesagt werden: **Der Einsatz von Zeigerinstrumenten zur Messung von Wechselspannungen oder -strömen in elektronischen Schaltungen ist ziemlich eng begrenzt.** Zur Untersuchung und Messung von Wechselstromvorgängen in elektronischen Schaltungen bedient man sich daher überwiegend des Oszilloskops.

1.3. Das Messen mit Oszilloskopen ¹⁾ (Oszillografen) ²⁾

Im Gegensatz zu Zeigerinstrumenten, die wegen ihrer mechanischen Trägheit bei Wechselstromvorgängen nur einen Mittelwert (Effektivwert) anzeigen können, arbeitet ein Oszilloskop praktisch trägheitslos. Zur „Anzeige“ dient bei einem Oszilloskop ein Elektronenstrahl, der beim Auftreffen auf einen Leuchtschirm einen Leuchtfleck hervorruft. Da Elektronen eine vernachlässigbar geringe Masse besitzen, kann man mit ihrer Hilfe äußerst schnelle elektrische Wechselvorgänge bildlich darstellen.

1.3.1. Grundsätzlicher Aufbau und Wirkungsweise eines Oszilloskops

1.3.1.1. Die Elektronenstrahlröhre

Kernstück eines Oszilloskops ist die Elektronenstrahlröhre, die nach ihrem Erfinder auch als Braunsche Röhre oder Braunsch's Rohr be-

¹⁾ Oszilloskop (griech.) = Schwingungsbetrachter

²⁾ Oszillograf (griech.) = Schwingungsschreiber

zeichnet wird. In ihr findet die Umkehrung des physikalischen Prinzips einer Fotozelle statt: Durch das Auftreffen eines im Vakuum stark beschleunigten Elektronenstrahls auf einen Leuchtschirm wird an der Auftreffstelle ein Leuchtfleck, also Licht, erzeugt. Wir wollen die Arbeitsweise einer Elektronenstrahlröhre anhand einer schematischen Darstellung erläutern.

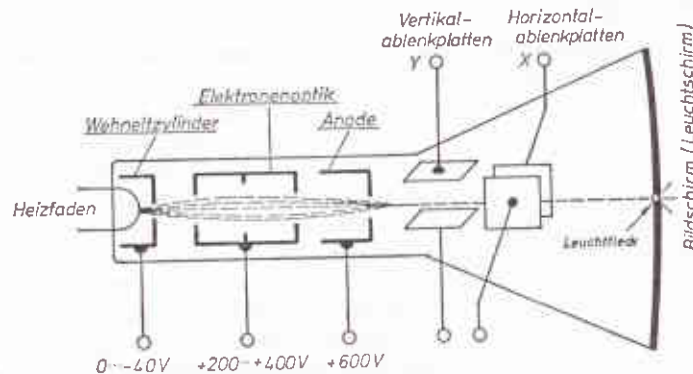


Abb. 1.15 — Schematischer Aufbau einer Elektronenstrahlröhre

Im gasleergepumpten Glaskolben der Röhre ist am sogenannten Halsende ein elektrisches Heizsystem angebracht. Der Heizfaden wird elektrisch erwärmt und sendet dabei freie Elektronen aus (Elektronenemission). Die negativ geladenen Elektronen werden von einer positiv vorgespannten Anode angezogen und auf der Strecke zwischen Heizelektrode und Anode stark beschleunigt. Da der Kolben gasleer ist, tritt den Elektronen kein Widerstand entgegen. Das Anodenblech besitzt eine kleine Öffnung, durch die die schnellen Elektronen hindurchfliegen und auf den Leuchtschirm der Röhre treffen. Der Leuchtschirm ist mit einem Stoff — z.B. Zinksulfid — beschichtet, der beim Auftreffen von Elektronen Licht erzeugt. Auf der Strecke zwischen der Heizung und der Anode liegen noch zwei Elektroden: Der Wehneltzylinder und die sogenannte Elektronenoptik. Man kann den Wehneltzylinder mit dem Steuergitter einer Elektronenröhre vergleichen. Durch Anlegen eines mehr oder weniger hohen negativen Potentials kann der Elektronenstrom in seiner Stärke und damit die Helligkeit des Leuchtflecks beeinflusst werden (Regler: INTENSITY). Da sich gleichnamige Ladungen abstoßen, ist der Elektronenstrom durch ein entsprechend hohes negatives Potential am Wehneltzylinder auch völlig auf Null zu drosseln. Der Wehneltzylinder besitzt wie die Anode eine nur kleine Öffnung, so daß nur ein feiner Elektronenstrahl austreten kann. Infolge geringer Streuungen würden aber die austretenden Elektronen nach einer gewissen Strecke nicht mehr einen feinen Strahl darstellen, sondern vielmehr ein Strahlenbündel mit zunehmend größerem Durchmesser ergeben. Um das zu verhindern, ist die Elektronenoptik da. Sie wirkt für den Elektronenstrahl wie eine Sammellinse und bündelt ihn so, daß er in Höhe des Leuchtschirms nur noch einen sehr feinen Strahl darstellt. Je feiner der Strahl, desto kleiner und schärfer ist der Leuchtfleck auf dem Leuchtschirm. Durch Verändern der Spannung an der Elektronenoptik läßt sich die Schärfe beeinflussen (Regler: FOCUS). Die Schärfeneinstellung wird geringfügig von der Helligkeitseinstellung beeinflusst.

In dem bis jetzt betrachteten Zustand entsteht lediglich ein winzig kleiner Leuchtfleck genau in der Mitte des Bildschirms. Der Elektronenstrahl kann nun mit Hilfe von zwei

Plattenpaaren, der Vertikal- und der Horizontalablenkung, durch Anlegen einer Spannung aus seiner Richtung abgelenkt werden. Aufgrund des Kraftwirkungsgesetzes der Elektrostatik werden die negativ geladenen Elektronen jeweils von der positiv geladenen Platte angezogen bzw. von der negativ geladenen Platte abgestoßen; man wendet also bei Oszilloskopen die elektrostatische Ablenkung an.

Im Gegensatz dazu wird der Elektronenstrahl in der Bildröhre eines Fernsehempfängers magnetisch mit Hilfe stromdurchflossener Spulen abgelenkt. Das ist bei der Fernsehbildröhre wegen des großen Bildschirms und der geringen Bildröhrenlänge notwendig, erfordert aber einen größeren Aufwand und ist für eine meßtechnische Auswertung im allgemeinen nicht genau genug.

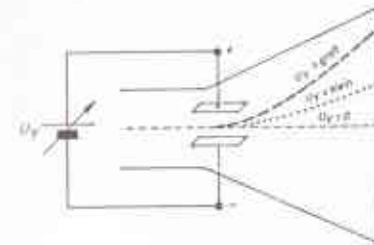


Abb. 1.16 — Einfluß der Höhe der Ablenkspannung auf den Elektronenstrahl

schlußbuchsen werden am Oszilloskop mit **Y-Eingang** bzw. **VERTICAL INPUT** und **X-Eingang** bzw. **HORIZONTAL INPUT** bezeichnet.

Soll auf dem Bildschirm ein Diagramm — z.B. der zeitliche Verlauf einer Wechselspannung — dargestellt werden, so muß der Elektronenstrahl in einer bestimmten Zeit von links nach rechts abgelenkt werden. Um z.B. den Verlauf einer Periode einer 50-Hz-Wechselspannung abbilden zu können, muß der Elektronenstrahl in 1/50 Sekunde von links nach rechts gewandert sein. Legen wir gleichzeitig an den Y-Eingang die Wechselspannung, so wird der Strahl dem jeweiligen Augenblickswert der Spannung entsprechend mehr oder weniger nach oben oder unten abgelenkt. Der durch den Elektronenstrahl hervorgerufene Leuchtfleck wandert in Form einer Sinuskurve von links nach rechts über den Bildschirm.

Damit man nun auf dem Bildschirm nicht nur den kleinen Leuchtfleck gerade an der jeweiligen Stelle erkennen kann, sondern den gesamten Verlauf seiner Bahn, muß das Bildschirmmaterial eine bestimmte Nachleuchtdauer aufweisen. Eine Unterschreitung der Nachleuchtdauer (langsame Vorgänge) hat ein Bildflimmern zur Folge. Je nach Verwendungszweck gibt es Elektronenstrahlröhren mit sehr kurzer bis zu extrem langer Nachleuchtdauer. Das auf dem Bildschirm eines Oszilloskops dargestellte Diagramm wird allgemein als **Oszillogramm** bezeichnet.

Zusammengefaßt kann gesagt werden:

In der Elektronenstrahlröhre wird durch das Auftreffen stark beschleunigter Elektronen auf einem Bildschirm ein Leuchtfleck erzeugt. Die Helligkeit läßt sich durch Verändern der negativen Vorspannung am Wehneltzylinder beeinflussen. Damit nur ein kleiner, scharfer Lichtpunkt entsteht, wird der Elektronenstrahl mit einer Elektronen-

optik gebündelt. Die Wirkung der Elektronenoptik ist durch Regeln der anliegenden Spannung beeinflussbar. Der Elektronenstrahl kann durch zwei Plattenpaare mit entsprechendem elektrischen Potential aus seiner Bahn abgelenkt werden:

Waagerechte Ablenkung = Horizontalablenkung = X-Ablenkung,

senkrechte Ablenkung = Vertikalablenkung = Y-Ablenkung.

1.3.1.2. Die Zeitablenkung des Elektronenstrahls

Zur Darstellung des zeitlichen Verlaufs einer Spannung muß der Elektronenstrahl in der Röhre in waagerechter Richtung jeweils von links nach rechts abgelenkt werden. Das geschieht durch Anlegen von Sägezahnimpulsen an die Horizontalablenkplatten. Die Sägezahnimpulse müssen einen absolut geradlinigen Flankenanstieg aufweisen, damit die Ablenkung in gleichen Zeitabschnitten zu gleichen Ablenkstrecken führt. Verfolgen wir einmal den Verlauf eines Sägezahnimpulses und seine Wirkung auf den Elektronenstrahl:

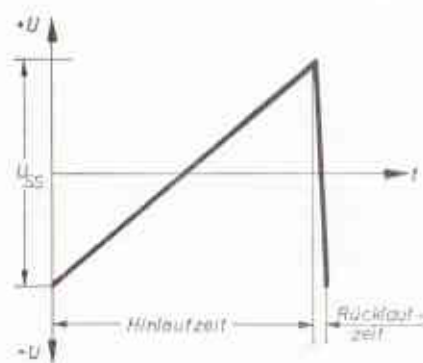


Abb. 1.17 — Zeitlicher Verlauf der zur Zeitablenkung dienenden Sägezahnspannung

Der Impuls beginnt mit dem negativen Höchstwert und verläuft im Laufe der Zeit linear über die Nulllinie hinweg zum positiven Höchstwert. Da es sich um einen gleichmäßigen Spannungsanstieg handelt, wandert der Elektronenstrahl infolge des elektrischen Feldes in gleichmäßiger Geschwindigkeit vom linken zum rechten Bildschirmrand. Bei der Spannung Null durchläuft der Elektronenstrahl die Bildschirmmitte (keine seitliche Ablenkung). Die Höhe der Sägezahnspannung läßt sich mit dem Regler HORIZONTAL AMPLITUDE einstellen.

Ebenso ist eine seitliche Verschiebung des Oszillogramms mit Hilfe des Reglers X-POSITION möglich. Ist der Elektronenstrahl mit Erreichen des Höchstwerts des Sägezahnimpulses am rechten Bildrand angelangt, so muß er in sehr kurzer Zeit an den Ausgangspunkt, also den linken Bildschirmrand, zurückgeführt werden. Die Ablenkspannung kippt daher sehr schnell vom positiven zum negativen Höchstwert; der dafür benötigte Zeitraum heißt Rücklaufzeit. Während der Rücklaufzeit springt der Leuchtfleck an den Ausgangspunkt. Nun

kann z.B. durch eine Kondensatorumladung ein nur für die Dauer der Rücklaufzeit wirkender negativer Spannungsstoß erzeugt werden. Führt man diesen negativen Spannungsstoß dem Wehneltzylinder der Elektronenstrahlröhre zu, so wird der Elektronenstrahl kurzzeitig unterbrochen. Diese Maßnahme findet Anwendung, um den Rücklauf des Elektronenstrahls auf dem Bildschirm unsichtbar zu machen.

Sollen auf dem Bildschirm periodische Vorgänge dargestellt werden — z.B. der Verlauf einer sinusförmigen Wechselspannung — so müssen die Folgefrequenz der Sägezahnimpulse und die Frequenz der zu untersuchenden Wechselspannung in einem genau ganzzahligen Verhältnis zueinander stehen. Um also z.B. eine Periode einer 50-Hz-Wechselspannung auf dem Bildschirm darzustellen, muß der Elektronenstrahl in 1/50 Sekunde vom linken zum rechten Bildrand abgelenkt werden, d.h., daß auch die Folgefrequenz der Ablenkimpulse 50 Hz betragen muß. Weicht die Folgefrequenz in unserem Beispiel auch nur geringfügig von 50 Hz ab, so wandert das Oszillogramm nach links oder rechts. Läßt man dagegen den Sägezahnimpuls in 1/25 Sekunde vom negativen zum positiven Höchstwert steigen, so erscheinen auf dem Bildschirm zwei Perioden der 50-Hz-Wechselspannung; denn in 1/25 Sekunde sind gerade zwei Perioden der Wechselspannung abgelaufen.

Die einfachen Beispiele lassen bereits erkennen, daß die Folgefrequenz der Ablenkspannung und die Frequenz der Meßspannung in einem ganzzahligen Verhältnis stehen müssen. Um das zu erreichen, werden je nach Art des Oszilloskops eines der folgenden oder auch beide Verfahren angewandt:

- a) Die Synchronisation des Kippgenerators und
- b) das Triggern eines Impulsgenerators mit Hilfe einer besonderen Triggerschaltung.

Bei der Synchronisation wird in einem besonderen, selbständig schwingenden Kippgenerator eine in ihrer Frequenz regelbare Sägezahnspannung erzeugt. Um den Gleichlauf zwischen Kippfrequenz und Frequenz der Meßspannung — also ein stehendes Schirmbild — zu erreichen, erhält der Kippgenerator einen kleinen Spannungsanteil von der Meßspannung und wird von deren Frequenz nachgesteuert. Treten jetzt im Verhältnis zwischen der Kippfrequenz und der Meßspannungsfrequenz Änderungen auf, so wird der Kippgenerator wieder auf die richtige Frequenz — nämlich ein ganzzahliges Vielfaches der Meßfrequenz — gezogen. Die Synchronisation ist vergleichbar mit der automatischen Scharfabbildung (AFC) eines UKW-Rundfunkempfängers.

Technisch eleganter und in modernen Oszilloskopen fast ausschließlich gebräuchlich ist das Triggern. Die Triggerschaltung schwingt im Gegensatz zu einem Kippgenerator nicht selbständig, sondern erzeugt nur dann einen einmaligen Sägezahnimpuls, wenn eine Steuerspan-

nung bestimmter Höhe angelegt wird. Als Steuerspannung dient meistens ein Anteil aus der Meßspannung. Zum Auslösen des Impulses genügt eine sehr geringe Steuerspannung, d.h. ein niedriger Schwellwert. Damit jedoch eine unerwünschte Auslösung durch Störanteile der Meßspannung vermieden werden kann, ist die Empfindlichkeit der Triggerschaltung durch einen Regler NIVEAU einstellbar. Erreicht die dem Y-Eingang zugeführte Meßspannung einen bestimmten, durch den Niveauregler eingestellten Schwellwert, so wird ein Sägezahnimpuls bestimmter Dauer ausgelöst. Dabei braucht die Meßspannung nicht einmal frequenzkonstant zu sein.

Mit einem Stufenschalter TIME BASE (Zeitbasis) kann die Dauer des Sägezahnimpulses eingestellt werden. Die Schalterstellungen sind in $\mu\text{s}/\text{cm}$ oder ms/cm geeicht, woraus sich ergibt, in welcher Zeit der Leuchtfleck in waagerechter Richtung um 1 cm abgelenkt wird. (Dabei müssen für Meßzwecke die anderen, ebenfalls die Bildbreite beeinflussenden Regler Horizontal-Amplitude und Vernier auf bestimmten Eichmarken stehen.) Aufgrund der Eichung der Zeitbasis-Einstellung ist es möglich, periodische Vorgänge sowohl niedriger als auch hoher Frequenz zu untersuchen oder wahlweise eine oder mehrere Perioden der Meßwechselspannung abzubilden und die Frequenz zu bestimmen.

Die Triggerschaltung kann aber auch selbstschwingend arbeiten — z.B. zur Darstellung von Gleichspannungen, die ja nur einen einzigen Triggerimpuls auslösen würden. Dazu ist eine besondere Schalterstellung AT (automatische Triggerung) vorhanden. Zur Untersuchung elektrischer Vorgänge ist manchmal auch die Steuerung der Triggerschaltung durch eine besondere Steuerspannung wünschenswert. Zu diesem Zweck kann der Steuereingang der Trigger- oder Synchronisationschaltung auf eine an der Frontplatte (extern) vorhandene Buchse geschaltet werden (Schalter: SYNCHRONISATION INTERN-EXTERN).

1.3.1.3. Die Y-Verstärkung

Die Höhe des Oszillogramms hängt von der Höhe der an die Vertikalablenkplatten angelegten Spannung ab. Damit ein Oszilloskop universell verwendbar ist, müssen mit ihm sowohl hohe als auch niedrige Spannungen nachgewiesen werden können. Darum liegt hinter der mit VERTICAL INPUT bezeichneten Buchse für den Y-Eingang ein als Stufenpotentiometer geschalteter Abschwächer VERTICAL AMPLITUDE und ein im Verstärkungsgrad umschaltbarer Verstärker. Dadurch ist die Einstellung der gewünschten Bildhöhe möglich. Als Verstärker dienen heute meistens galvanisch gekoppelte Transi-

storverstärker. Vor den Y-Eingang kann mit Hilfe eines Schalters DC—AC ein Kondensator geschaltet werden, der dem Oszilloskop ausschließlich den Wechselspannungsanteil der Meßspannung zuführt. In Stellung AC ist dieser Kondensator vor den Eingang geschaltet.

Um anhand der Ablenkweite auf dem Bildschirm die Höhe der Spannung bestimmen zu können, wird für Oszilloskope der **Ablenkkoefizient** angegeben. Soweit nur eine Angabe über den Ablenkkoefizient gemacht wird, handelt es sich um den Wert bei größtmöglicher Verstärkung des Vertikalverstärkers. **Der Ablenkkoefizient gibt an, wie hoch die Spannung am Eingang sein muß, damit der Leuchtfleck auf dem Bildschirm um 1 cm abgelenkt wird.** Seine Einheit ist entsprechend V/cm bzw. mV/cm . Er wird jedoch — wie schon angedeutet — durch den Eingangsabschwächer und die Verstärkungsregelung beeinflusst. Deshalb sind sowohl Eingangsabschwächer als auch der Stufenschalter für die Verstärkung geeicht.

Zur Einstellung der Bildlage in senkrechter Richtung — z.B. auf die Bezugslinie eines vor dem Bildschirm liegenden Rasters — dient ein Regler Y-POSITION.

Da die Ablenkplatten für die Vertikalablenkung vom Bildschirm der Elektronenstrahlröhre weiter entfernt liegen, ist der Y-Eingang wesentlich empfindlicher als der X-Eingang eines Oszilloskops.

1.3.1.4. Blockschaltbild und Bedienplatte

Zur besseren Übersicht über das Zusammenwirken der einzelnen Baugruppen eines Oszilloskops ist in Abb. 1.18 das Blockschaltbild eines einfachen Oszilloskops dargestellt; die Abbildung enthält nur die wichtigsten Baugruppen bzw. Teile.

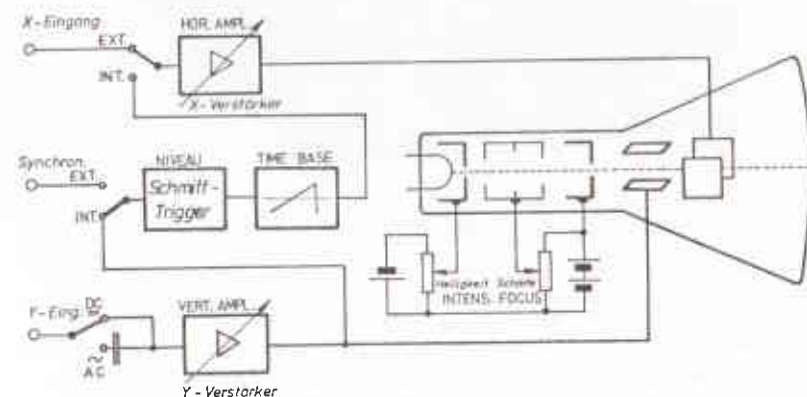


Abb. 1.18 — Blockschaltbild eines Oszilloskops

Alle Regler, Schalter und Anschlußbuchsen sind normalerweise auf der Frontplatte, der sogenannten Bedienplatte, angeordnet. Die Ansicht eines auch für Ausbildungszwecke gebräuchlichen Oszilloskops ist in Abb. 1.19 wiedergegeben. Daneben besteht bei einigen Typen die Möglichkeit, Meßspannungen und Ablenkspannungen direkt an die Ablenkplatten zu legen. Dabei werden z.B. Einflüsse von Schaltungskapazitäten vermieden. Die Anschlußbuchsen der Ablenkplatten liegen meistens unter dem hinteren Gehäuseabschluß.

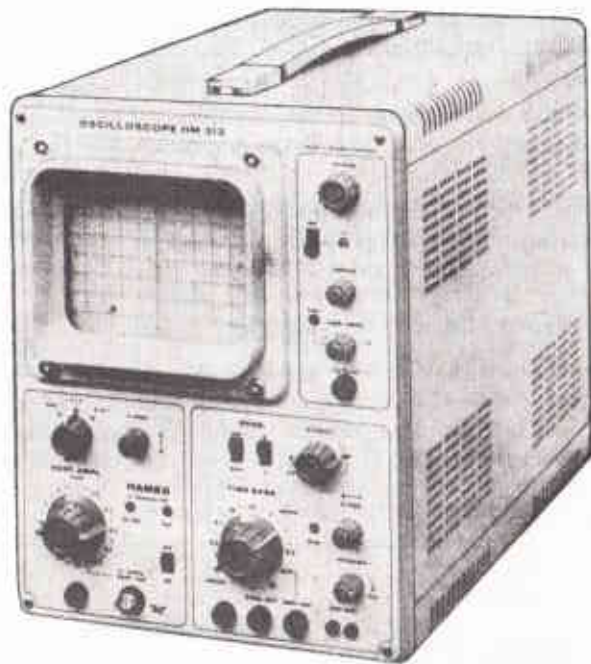


Abb. 1.19 — Ansicht eines Triggeroszilloskops

Die Beschriftung der Bedienelemente wird heute allgemein — auch bei Geräten deutscher Herstellung — in englisch vorgenommen. Zur Vermeidung von Bedienungsfehlern sollten Zweck und Bedeutung der Bedienelemente sicher beherrscht werden. Eine Übersicht mit kurzen Erläuterungen finden Sie auf Seite 30.

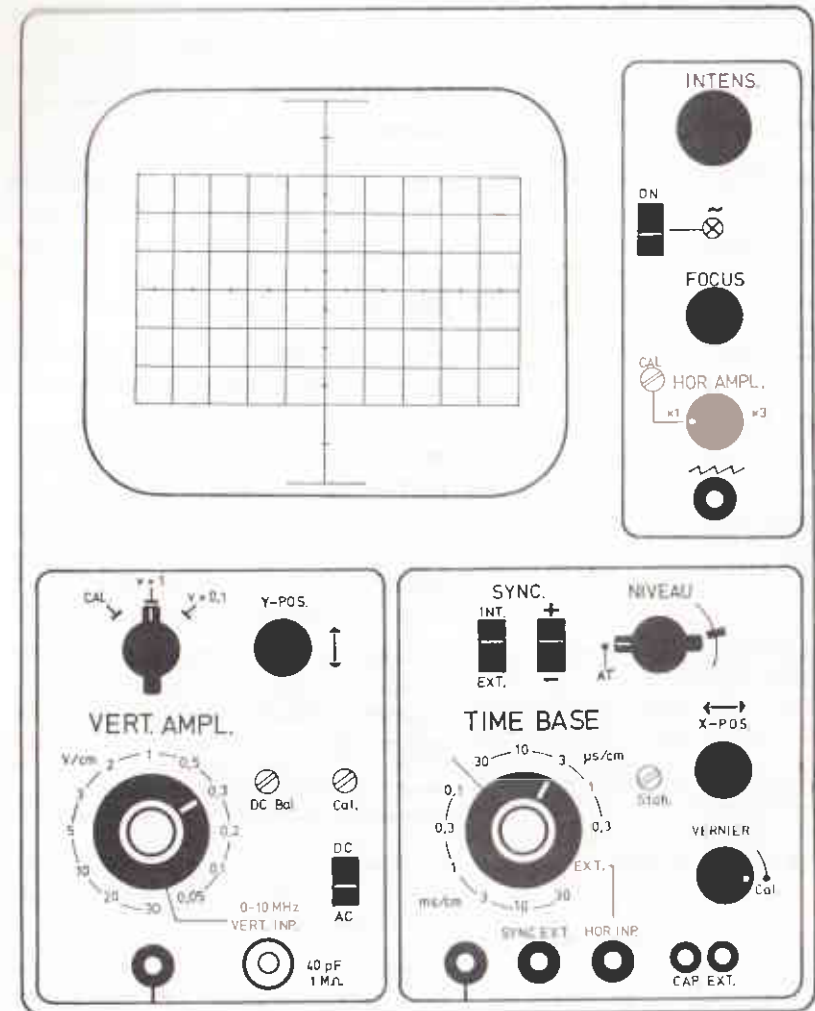


Abb. 1.20 — Frontplatte eines Triggeroszilloskops

Zusammenstellung der wichtigsten Bedienungselemente, deren Bezeichnung und Bedeutung

POWER ON—OFF	= Netzschalter (ON=Ein, OFF=Aus)
INTENSITY	= Helligkeitsregler
FOCUS	= Bildschärferegler
HORIZONTAL AMPLITUDE	= Regler mit markierter Ruhestellung zur Beeinflussung der Bildbreite
TIME BASE	= Zeitbasis-Stufenschalter zur Einstellung der Länge (Dauer) eines Horizontal-Sägezahnimpulses
VERNIER	= Regler mit markierter Ruhestellung zur Feineinstellung der Zeitbasis
NIVEAU	= Regler zur Einstellung des Schwellwerts zum Triggern der Horizontalablenkung
SYNC. INTERN—EXTERN	= Umschalter zur wahlweisen Synchronisation durch die Meßspannung selbst (INT.) oder durch eine besondere Steuerspannung (EXT.)
SYNCHRONISATION + —	= Umschalter zur wahlweisen Synchronisation durch die positive oder negative Anstiegsflanke der Meßspannung
VERTICAL AMPLITUDE	= geeichter Stufenschalter für den Eingangsabschwächer (Y-Verstärker)
$v \times 1$ $v \times 0,1$	= Umschalter für den Verstärkungsgrad des Y-Verstärkers (größte Verstärkung bei $v \times 0,1$)
DC — AC	= Eingangswahlschalter für Gleichstrom (DC) oder Wechselstrom (AC) des Y-Verstärkers
Y-POSITION	= Regler zur senkrechten Verschiebung des Oszillogramms
X-POSITION	= Regler zur waagerechten Verschiebung des Oszillogramms

1.3.1.5. Allgemeine Hinweise zum Gebrauch eines Oszilloskops

Das Ein- und Ausschalten des Geräts soll nur mit Hilfe des eingebauten Netzschalters — und nicht etwa durch Herausziehen des Netzsteckers — vorgenommen werden, da sonst Brennflecke auf dem Bildschirm entstehen können. Die Elektronenstrahlröhre wird dadurch unbrauchbar. Bei einfallendem Tageslicht ist der Kontrast des Oszillogramms meistens zu gering. Eine Kontrasterhöhung ist mit einem Kontrastfilter in der Eigenfarbe des Leuchtflecks möglich. Der Y-Eingang ist bei der Einstellung auf größte Verstärkung äußerst empfindlich. Schon ein kurzes angeschlossenes Drahtstück wirkt wie eine Antenne. So zeigt sich bereits, ohne daß eine Meßspannung angelegt worden ist, ein Oszillogramm, das von Störfeldern verursacht wird (z.B. Rundfunksender). Zur Messung sehr kleiner Spannungen muß daher auf eine besonders gute Erdung des Geräts geachtet werden. Vielfach müssen auch die Zuleitungen — vor allem bei hohem Innenwiderstand der Steuerspannungsquelle — abgeschirmt sein.

Bei der Inbetriebnahme eines Oszilloskops ist unbedingt darauf zu achten, daß der Helligkeitsregler INTENSITY zurückgeregelt ist. Solange nämlich keine Ablenkspannung anliegt oder der Regler NIVEAU nicht auf AT steht, bleibt der Leuchtfleck in der Mitte des Bildschirms stehen. Durch eine bei starker Helligkeit übermäßige Erwärmung wird die Leuchtschicht an dieser Stelle verbrannt und somit unbrauchbar.

Infolge des hohen Eingangswiderstands ergibt sich in Verbindung mit der im AC-Eingang liegenden Kapazität eine verhältnismäßig große Zeitkonstante. Es ist daher nicht ungewöhnlich, wenn das Oszillogramm nach dem Umschalten des DC-AC-Schalters erst nach einer gewissen Zeit seine endgültige Höhenlage einnimmt.

Die verdeckt angebrachten Einsteller, z.B. „DC-Balance“ oder „Cal.“ oder „Stab.“ sollten auf keinen Fall betätigt werden. Sie dienen zur genauen Eichung des Oszilloskops und können im Bedarfsfall nur mit genauen Meßgeräten eingestellt werden.

Gegenüber Zeigermeßinstrumenten ist der Eingangswiderstand eines Oszilloskops erheblich hochohmiger und beträgt ein Megohm bis 11 Megohm. Aus diesem Grunde **eignen sich Oszilloskope auch nur zur Messung oder Darstellung von Spannungen. Soll die Stromstärke Meßgröße sein, so ist immer die indirekte Strommessung anzuwenden.** Die indirekte Strommessung wurde bereits im Abschn. 1.2 „Messungen mit Zeigerinstrumenten“ beschrieben. Bei niedrigen Frequenzen (≤ 25 Hz) ist wegen der begrenzten Nachleuchtdauer des Bildschirms kein flimmerfreies Bild mehr zu erreichen.

1.3.2. Der elektronische Schalter als Zusatzgerät

Ein elektronischer Schalter als Zusatzgerät zum Oszilloskop erweitert dessen Anwendungsbereich erheblich. Durch seine Verwendung können gleichzeitig zwei Oszillogramme auf dem Bildschirm sichtbar gemacht werden. Die Arbeitsweise des elektronischen Schalters beruht darauf, daß er dem Y-Eingang des Oszilloskops in schneller Folge abwechselnd zwei verschiedene Meßspannungen zuschaltet. Mechanische Schalter arbeiten wegen ihrer mechanischen Trägheit viel zu langsam; darum wird die Umschaltung mit Hilfe von Röhren oder Transistoren als Schalter vorgenommen. Um Flimmererscheinungen zu vermeiden, sollte die Umschaltfrequenz mindestens 25 Hz betragen; sie ist regel- oder umschaltbar. Das Arbeitsprinzip eines elektronischen Schalters geht aus Abb. 1.21 hervor. Zur Veranschaulichung ist dabei anstelle des elektronischen Umschalters ein Umschaltkontakt dargestellt. Am elektronischen Schalter läßt sich für jeden Kanal getrennt die Spannungshöhe regeln und eine vertikale Bildverschiebung vornehmen.

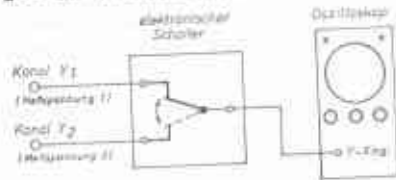


Abb. 1.21 — Prinzipschaltung eines elektronischen Schalters



Abb. 1.22 — Oszillogramm zweier um 90° phasenverschobener Wechselspannungen

In Abb. 1.22 ist das mit Hilfe eines elektronischen Schalters dargestellte Oszillogramm zweier um 90° phasenverschobener 50-Hz-Wechselspannungen wiedergegeben. Die Umschaltfrequenz beträgt dabei 800 Hz.

Der elektronische Schalter läßt z.B. folgende Untersuchungen und Messungen an einem einfachen Oszilloskop zu:

1. Nachweis und Messung der **Phasenverschiebung** zwischen zwei Spannungen, zwei Strömen oder zwischen Spannung und Strom.
2. Darstellung der **Phasenlage** von Eingangs- und Ausgangsspannung an einem Vierpol, z.B. einem Verstärker.
3. **Vergleich der Kurvenformen** von Eingangs- und Ausgangsspannung an einem Vierpol, z.B. Verstärker oder Trigger, um erwünschte oder unerwünschte Verformungen des Signals nachzuweisen. Durch vertikale Verschiebung lassen sich zwei Oszillogramme auch untereinander darstellen.
4. Der elektronische Schalter läßt sich als **Rechteckgenerator** verwenden.

Für die Einstellung der Umschaltfrequenz gelten folgende Richtlinien:

- a) Um Bildflimmern zu vermeiden, sollte die Umschaltfrequenz nicht kleiner als 25 Hz sein.
- b) Eine Synchronisierung von Meßspannungsfrequenz und Umschaltfrequenz ist im allgemeinen nicht notwendig.
- c) Je höher die Umschaltfrequenz ist, desto genauer wird der Verlauf der beiden Meßspannungen abgebildet. Die höchstmögliche Umschaltfrequenz wird durch die größte einstellbare Frequenz am elektronischen Schalter und das Frequenzverhalten des Y-Eingangs am Oszilloskop begrenzt.

Beträgt z.B. die Frequenz der Meßspannungen 50 und die Umschaltfrequenz 500 Hz, so besteht das Oszillogramm jedes Kanals aus jeweils 10 kurzen Strichabschnitten.

- d) Liegt die Frequenz der Meßwechselspannungen hoch — etwa 300 Hz und größer —, so kann mit einer kleineren Umschaltfrequenz die Darstellung jeweils einer oder mehrerer vollständiger Perioden erreicht werden. Dabei wird abwechselnd im Rhythmus der Umschaltfrequenz einmal das Oszillogramm für Kanal 1 und zum anderen für Kanal 2 geschrieben. Bei entsprechender vertikaler Verschiebung erscheinen die Bilder untereinander.

1.3.3. Einfache Messungen mit dem Oszilloskop

Es empfiehlt sich, den Umgang mit dem Oszilloskop anhand einfacher Beobachtungen und Messungen zu üben. Auf die Auswertung von Oszillogrammen an elektronischen Bauelementen und Schaltungen soll daher erst im Zusammenhang mit den entsprechenden Abschnitten eingegangen werden.

1.3.3.1. Gleichspannungsmessungen

Soll ein Oszilloskop nicht nur zur Beobachtung, sondern auch zur Messung elektrischer Vorgänge dienen, so sollte vor dem Bildschirm eine Rasterscheibe liegen. Die Rasterscheibe ermöglicht die Festlegung einer Bezugslinie oder eines Bezugspunkts. Zur Gleichspannungsmessung wird der Leuchtfleck zunächst ohne Meßspannung mit Hilfe der X- und Y-Positionsregler auf den Mittelpunkt (Koordinaten-

schnittpunkt) der Rasterscheibe eingestellt und anschließend die Horizontalablenkung auf automatische Triggerung (AT) geschaltet. Der Leuchtstrich muß dabei genau auf der waagerechten Bezugsachse des Rasters liegen.

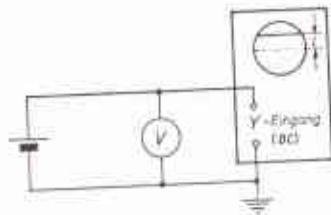


Abb. 1.23 — Gleichspannungsmessung mit dem Oszilloskop

Nun wird die Meßspannung an den Y-Eingang (VERTICAL INPUT) gelegt und mit Hilfe der beiden Schalter VERTICAL AMPLITUDE eine Ablenkweite von etwa 3 cm eingestellt.

Die Ablesegenauigkeit läßt sich bei Gleichspannungsmessungen noch dadurch erhöhen, daß man als Bezugslinie die untere waagerechte Linie des Rasters wählt und damit eine größere Ablenkung ermöglicht.

Die Höhe der Spannung ergibt sich dann aus dem Produkt von gemessener senkrechter Ablenkweite l , eingestelltem Ablenkkoeffizient AK des Schalters VERTICAL AMPLITUDE und dem eingestellten Faktor v des Verstärkungsschalters.

$$U = l \cdot AK \cdot v$$

Beispiel:

Beim Anlegen einer Gleichspannung an den Y-Eingang wird auf dem Bildschirm eine Ablenkweite von 2,6 cm gemessen. Der Stufenschalter VERTICAL AMPLITUDE steht auf 3 V/cm, der Stufenschalter des Y-Verstärkers auf $v \times 0,1$. Wie hoch ist die Gleichspannung?

Gegeben: $l = 2,6 \text{ cm}$; $AK = 3 \text{ V/cm}$; $v = 0,1$

Gesucht: U

Lösung: $U = l \cdot AK \cdot v$

$$U = 2,6 \text{ cm} \cdot 3 \frac{\text{V}}{\text{cm}} \cdot 0,1 = 0,78 \text{ V}$$

Die Eingangsspannung beträgt 0,78 Volt.

Das rechnerische Ergebnis muß mit dem Anzeigewert eines Zeigerinstruments übereinstimmen.

1.3.3.2. Gleichstrommessung

Wegen des hohen Eingangswiderstands lassen sich mit einem Oszilloskop keine direkten Strommessungen durchführen. Man muß daher die indirekte Meßmethode für Ströme anwenden. Dazu wird in den Meßstromkreis ein möglichst kleiner Widerstand geschaltet. Der Spannungsabfall an diesem Widerstand ist dann ein Maß für die über ihn fließende Stromstärke ($I = U/R$). Wie klein der Widerstand gewählt werden kann, richtet sich nach dem kleinsten Ablenkkoeffizient und der zu erwartenden Meßstromstärke; denn mit Rücksicht auf eine gute Ablesegenauigkeit sollte das Oszillogramm eine bestimmte Mindesthöhe von etwa $\frac{1}{3}$ bis $\frac{1}{4}$ der Bildschirmhöhe nicht unterschreiten.

Beispiel:

Der kleinste einstellbare Ablenkkoeffizient eines Oszilloskops beträgt 0,05 V/cm. Dabei wird der Schalter für die Verstärkung auf $v \times 0,1$ gestellt. Im Meßstromkreis wird mit einer Stromstärke von mindestens etwa 1 mA gerechnet. Wie groß muß der Meßwiderstand sein, damit eine Mindestablenkweite von 2 cm auftritt?

Gegeben: $AK = 0,05 \text{ V/cm}$; $v = 0,1$; $l = 2 \text{ cm}$; $I = 1 \text{ mA}$

Gesucht: R

Lösung: $R = \frac{U}{I}$

$$U = l \cdot AK \cdot v$$

$$U = 2 \text{ cm} \cdot 0,05 \frac{\text{V}}{\text{cm}} \cdot 0,1 = 0,010 \text{ V}$$

$$R = \frac{0,010 \text{ V}}{0,001 \text{ A}} = 10 \Omega$$

Der Meßwiderstand muß einen Wert von 10 Ohm haben.

Bei der Wahl des Widerstandswerts sollte wegen der einfacheren Stromberechnung immer nach Möglichkeit ein runder Dezimalwert, wie z.B. 0,1, 10, 100 usw. bevorzugt werden.

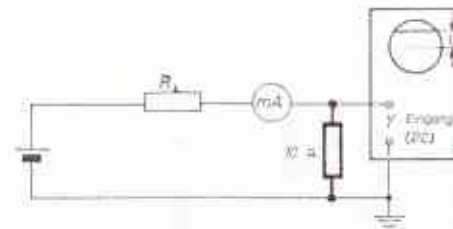


Abb. 1.24 — Gleichstrommessung mit dem Oszilloskop

Da die indirekte Strommessung eigentlich eine Spannungsmessung ist, unterscheidet sich das Meßverfahren nicht von der Gleichspannungsmessung. Aus der Ablenkweite läßt sich die Stromstärke bestimmen.

Beispiel:

Bei einem eingestellten Ablenkkoeffizient $0,05 \text{ V/cm}$ bei $v = 0,1$ wird am Bildschirm eine Ablenkweite von $2,8 \text{ cm}$ gemessen. Der Meßwiderstand hat einen Wert von 10 Ohm . Wie groß ist die Stromstärke im Meßstromkreis?

Gegeben: $AK = 0,05 \text{ V/cm}; v = 0,1; l = 2,8 \text{ cm}; R = 10 \Omega$

Gesucht: I

Lösung:

$$I = \frac{U}{R}$$

$$U = I \cdot AK \cdot v$$

$$U = 2,8 \text{ cm} \cdot 0,05 \frac{\text{V}}{\text{cm}} \cdot 0,1 = 0,014 \text{ V}$$

$$I = \frac{0,014 \text{ V}}{10 \frac{\text{V}}{\text{A}}} = 0,0014 \text{ A} = \underline{\underline{1,4 \text{ mA}}}$$

Die Stromstärke beträgt $1,4 \text{ Milliampere}$.

Der errechnete Stromwert muß mit dem von einem Zeigerinstrument angezeigten Wert übereinstimmen.

1.3.3.3. Wechselspannungsmessung

Die Messung der Höhendifferenz am Wechselspannungs-Oszillogramm wird dadurch erschwert, daß die positiven und negativen Höchstwerte seitlich (eigentlich zeitlich) gegeneinander versetzt sind. Genaue Messungen der Ablenkweiten zwischen positiven und negativen Maximalwerten sind daher nur mit Hilfe fein unterteilter Rasterscheiben möglich. Man kann sich dabei allerdings eines einfachen Kniffs bedienen. Schaltet man nämlich die Horizontalablenkung aus (Zeitbasisschalter

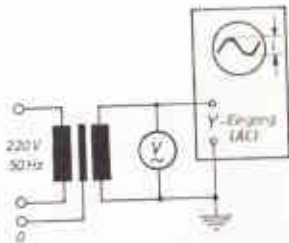


Abb. 1.25 — Wechselspannungsmessung mit dem Oszilloskop

auf EXT.), so erscheint die Wechselspannung auf dem Bildschirm als ein senkrechter Strich, dessen Länge einfach ausgemessen werden kann. Dieses Verfahren sollte jedoch nur angewandt werden, nachdem die Kurvenform der Wechselspannung durch Beobachten des zeitabgelenkten Oszillogramms festgestellt worden ist.

Vergleicht man das Ergebnis anhand des Oszillogramms mit der Anzeige eines Zeigerinstruments, so zeigt sich, daß die Werte nicht übereinstimmen. Da der Elektronenstrahl keine Trägheit besitzt, folgt er der Wirkung der Ablenkspannung bis zum jeweiligen Spitzenwert. Anhand des Oszillogramms ergibt sich also immer der Spitze-

Spitze-Wert. Bei sinusförmiger Wechselspannung läßt sich aus diesem Wert der Effektivwert berechnen:

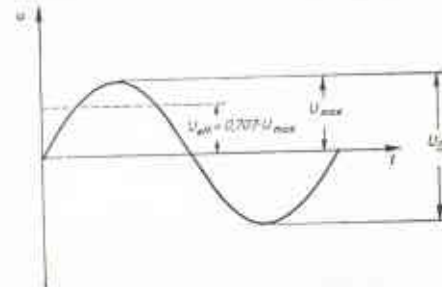


Abb. 1.26 — Charakteristische Größen einer sinusförmigen Wechselspannung

$$U_{\text{eff}} = 0,707 \cdot U_{\text{max}}$$

$$U_{\text{max}} = \frac{U_{\text{SS}}}{2}$$

$$U_{\text{eff}} = \frac{0,707}{2} \cdot U_{\text{SS}}$$

$$U_{\text{eff}} = 0,3535 \cdot U_{\text{SS}}$$

Beispiel:

Auf dem Bildschirm eines Oszilloskops wird zwischen positivem und negativem Höchstwert einer sinusförmigen Wechselspannung eine Höhendifferenz von $5,5 \text{ cm}$ gemessen. Die beiden Schalter für VERTICAL AMPLITUDE stehen auf 2 V/cm und $v \times 1$. Wie hoch ist der Effektivwert der Wechselspannung?

Gegeben: $l = 5,5 \text{ cm}; AK = 2 \text{ V/cm}; v = 1$

Gesucht: U_{eff}

Lösung:

$$U_{\text{eff}} = 0,3535 \cdot U_{\text{SS}}$$

$$U_{\text{SS}} = l \cdot AK \cdot v$$

$$U_{\text{SS}} = 5,5 \text{ cm} \cdot 2 \frac{\text{V}}{\text{cm}} \cdot 1 = 11 \text{ V}$$

$$U_{\text{eff}} = 0,3535 \cdot 11 \text{ V} = \underline{\underline{3,89 \text{ V}}}$$

Der Effektivwert der Wechselspannung beträgt $3,89 \text{ Volt}$.

Zur Beurteilung von Wechselstromvorgängen ist es oft vorteilhaft, nur einen bestimmten Teil einer Periode oder auch mehrere Perioden nebeneinander auf dem Bildschirm zu betrachten. Mit Hilfe des Zeitbasisschalters sind diese Möglichkeiten gegeben. Der Einfluß unterschiedlicher Zeitbasiseinstellungen auf das Oszillogramm eines

periodischen Vorgangs soll am Beispiel einer 50-Hz-Wechselspannung erläutert werden.

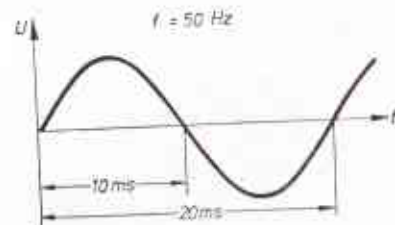


Abb. 1.27 — Zeit-Spannungs-Diagramm einer 50-Hz-Wechselspannung

Bei der 50-Hz-Wechselspannung dauert eine Periode $1/50$ Sekunde = 20 ms. Soll eine Periode auf einer Rasterbreite von 10 cm wiedergegeben werden, so müßte der Elektronenstrahl in 20 ms um 10 cm in waagerechter Richtung abgelenkt werden. Der zugehörige Zeitbasiswert beträgt demnach $20 \text{ ms}/10 \text{ cm} = 2 \text{ ms/cm}$. Da bei unserem Oszilloskop dieser Wert nicht einstellbar ist, wird der nächstgrößere Wert 3 ms/cm gewählt. Das bedeutet, daß es insgesamt 30 ms dauert, bis der Leuchtfleck um 10 cm nach rechts abgelenkt worden ist. Dabei ergibt sich das folgende Oszillogramm.

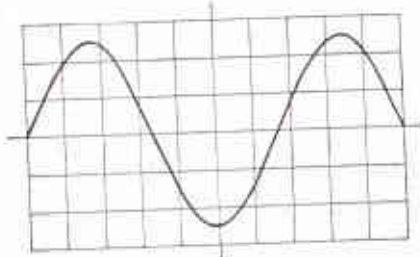


Abb. 1.28 — Oszillogramm einer 50-Hz-Wechselspannung bei einer Zeitbasis von 3 ms/cm

Soll nur eine Halbperiode abgebildet werden, so muß die Ablenkung in kürzerer Zeit erfolgen, die Zeitbasis also verringert werden. Gewählt wird der Zeitbasiswert 1 ms/cm; d.h., für 10 cm werden 10 ms benötigt. Diese Einstellung hat eine Lupenwirkung.

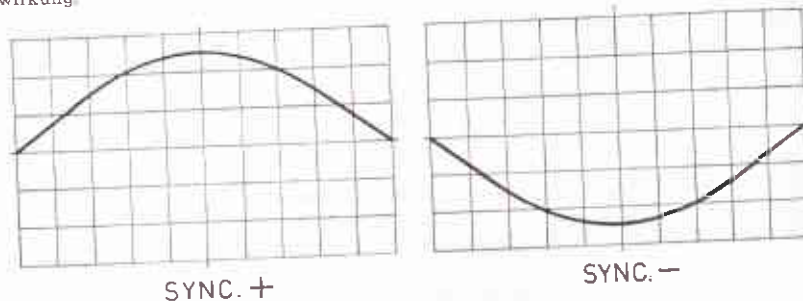


Abb. 1.29 — Oszillogramme einer 50-Hz-Wechselspannung bei einer Zeitbasis von 1 ms/cm

Durch Umschalten des Schalter SYNC. von + auf - läßt sich auch wahlweise die negative Halbwelle oszillografieren.

Wird die Zeitbasis erhöht, so braucht der Leuchtfleck längere Zeit, um von links nach rechts über den Bildschirm zu wandern. Bei einer Zeitbasis von 10 ms/cm würden für 10 cm Ablenkweite insgesamt 100 ms Zeit benötigt. In dieser Zeit sind aber 5 Perioden der 50-Hz-Wechselspannung abgelaufen.

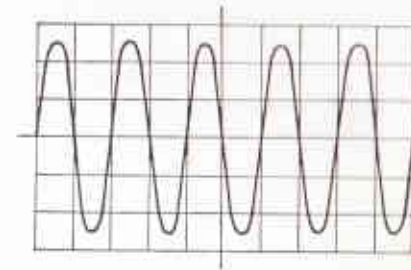


Abb. 1.30 — Oszillogramm einer 50-Hz-Wechselspannung bei einer Zeitbasis von 10 ms/cm

1.3.3.4. Wechselstrommessung

Auch zur Wechselstrommessung ist die indirekte Strommeßmethode anzuwenden. Bei der rechnerischen Auswertung der Messung ist unbedingt zu berücksichtigen, daß die gemessene Ablenkweite immer den Spitze-Spitze-Wert des Spannungsabfalls am Meßwiderstand ergibt. Soll für sinusförmigen Wechselstrom der Effektivwert bestimmt werden, so ist — wie bereits abgeleitet — mit dem Faktor 0,3535 zu multiplizieren.

1.3.3.5. Darstellung von Kennlinien

Da in der Elektronik überwiegend Bauteile verwendet werden, deren Kennlinien nicht formelmäßig zu erfassen sind, ist man auf die meßtechnische Kennlinienaufnahme angewiesen. Mit Zeigerinstrumenten müssen Kennlinien punktwise aufgenommen werden, was meistens sehr zeitraubend ist. Ein Oszilloskop ermöglicht dagegen eine vollständige Kennliniendarstellung.

In der Kennlinie für ein elektronisches Bauteil wird normalerweise die Abhängigkeit der Stromstärke von der Höhe der anliegenden Spannung dargestellt. Bei der Kennlinienaufnahme wird die Spannung stufenweise erhöht und zu jedem eingestellten Spannungswert die Stromstärke gemessen. Die Spannungswerte werden dann nach einem festgelegten Maßstab auf der waagerechten Achse, die zugehörigen Stromwerte ebenfalls maßstabsgerecht senkrecht dazu aufgetragen. Die Verbindung aller so erhaltenen Punkte ergibt dann die Kennlinie.

Da man den Elektronenstrahl in einer Elektronenstrahlröhre sowohl waagrecht als auch senkrecht ablenken kann, bietet sich ein Oszilloskop zur Darstellung von Diagrammen — also auch Kennlinien — geradezu an. Bei Spannungs-Strom-Kennlinien muß nur berücksichtigt werden, daß die Stromstärke beim Oszilloskop nur indirekt angezeigt werden kann. In Abb. 1.31 ist die Grundschaltung zur Kennlinienaufnahme eines ohmschen Widerstands wiedergegeben.

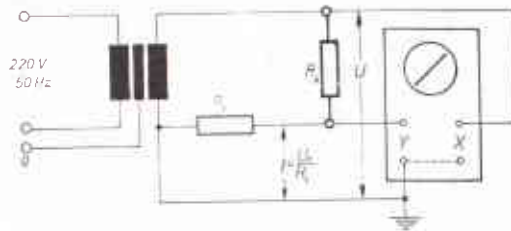


Abb. 1.31 — Schaltung zur Kennlinienaufnahme

Die Meßspannung wird dabei dem **Horizontaleingang** (Zeitbasis-Schalter auf EXT.) zugeführt. Der Spannungsabfall am niederohmigen Widerstand R_1 ist ein Maß für die **Stromstärke** und liegt am **Vertikaleingang**. Wird nun die Spannung an der Meßschaltung sehr schnell — z.B. 50-Hz-Wechselspannung — geändert, so zeigt sich auf dem Bildschirm die Kennlinie des Widerstands R_2 . Der Abbildungsmaßstab läßt sich waagrecht durch den Regler HORIZONTAL AMPLITUDE und senkrecht durch den Schalter VERTICAL AMPLITUDE wunschgemäß einstellen.

Hinweis: Die Meßschaltung enthält einen Fehler. Die Spannung U wird nicht nur an R_2 , sondern an der Reihenschaltung $R_2 - R_1$ abgegriffen. Das ist notwendig, weil die „unteren“ Buchsen des Horizontal- und Vertikaleingangs galvanisch miteinander verbunden sind (Masse). Würde der Verbindungspunkt $R_1 - R_2$ der Meßschaltung an den Masseanschluß des Oszilloskops gelegt, ergäbe sich ein um 180° gedrehtes Oszillogramm, das der üblichen Kennliniendarstellung nicht entspräche. Der Fehler in der „seltenrichtigen“ Meßschaltung läßt sich um so geringer halten, je kleiner R_1 im Verhältnis zu R_2 gemacht werden kann.

Nach dem hier angegebenen Verfahren können die Kennlinien von Halbleiter-Bauelementen wie Dioden, Z-Dioden, VDR-Widerständen usw. dargestellt werden. Die Bauelemente werden dazu anstelle des Widerstands R_2 in die Schaltung eingefügt.

1.3.3.6. Frequenzmessung mit dem Oszilloskop

Mit Hilfe eines Oszilloskops lassen sich Frequenzmessungen durchführen, wenn der Zeitbasis-Schalter geeicht und die Regler für Hori-

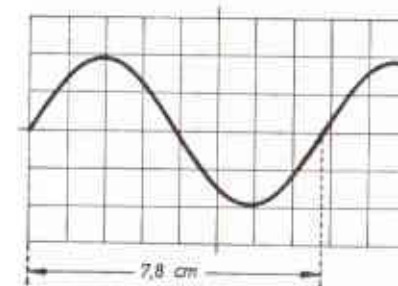
zontalablenkung und Zeitbasis auf Eichmarken stehen. **Die eingestellte Zeitbasis gibt an, in welcher Zeit der Leuchtfleck auf dem Bildschirm um 1 cm nach rechts abgelenkt wird.** Zur Frequenzmessung sollte grundsätzlich ein Raster vor dem Bildschirm liegen.

Zur Frequenzbestimmung legt man die Wechselspannung unbekannter Frequenz an den Vertikaleingang und stellt Bildhöhe und Zeitbasis so ein, daß die Bildschirmfläche durch eine Periode der Wechselspannung gut ausgenutzt wird. Mit Hilfe der waagrecht gemessenen Strecke zwischen Anfang und Ende einer Periode und dem Zeitbasiswert t_B lassen sich dann die Periodendauer T und die Frequenz f berechnen:

$$T = l \cdot t_B$$

$$f = \frac{1}{T}$$

Beispiel:



Auf dem Bildschirm eines Oszilloskops wird gemäß nebenstehender Abbildung zwischen Anfang und Ende einer Periode eine Strecke von 7,8 cm gemessen. Der Zeitbasis-Schalter steht auf $3 \mu\text{s}/\text{cm}$. Wie hoch ist die Frequenz?

Abb. 1.32 — Oszillogramm zur Frequenzbestimmung

Gegeben: $l = 7,8 \text{ cm}$; $t_B = 3 \mu\text{s}/\text{cm}$

Gesucht: f

Lösung: $f = \frac{1}{T}$

$$T = l \cdot t_B$$

$$T = 7,8 \text{ cm} \cdot 3 \frac{\mu\text{s}}{\text{cm}} = 23,4 \mu\text{s} = 23,4 \cdot 10^{-6} \text{ s}$$

$$f = \frac{1}{23,4 \cdot 10^{-6} \text{ s}} = \frac{10^6}{23,4} \frac{1}{\text{s}} = \underline{\underline{42700 \text{ Hz}}}$$

Die Frequenz der Wechselspannung beträgt 42,7 kHz.

2. Halbleiter

2.1. Was ist ein Halbleiter?

Bei der Einteilung der Stoffe in bezug auf ihre Leitfähigkeit für den elektrischen Strom unterscheidet man im allgemeinen zwischen

guten Leitern, z.B. Kupfer, Silber,

Halbleitern, z.B. Selen, Germanium, Silizium und

Nichtleitern, z.B. Glas, Porzellan.

Dabei werden die Halbleiter als Stoffe eingestuft, deren elektrische Leitfähigkeit zwischen der eines guten und der eines schlechten Leiters (Isolators) liegt. Die Grenze zwischen den drei Gruppen ist bei dieser Einteilung jedoch sehr ungenau.

Die Halbleiter unterscheiden sich von den guten Leitern viel eindeutiger durch ihr Temperaturverhalten. Während der Widerstand der guten Leiter mit steigender Temperatur zunimmt, wird der Widerstand der Halbleiter bei zunehmender Erwärmung immer geringer. Umgekehrt fällt der Widerstand der guten Leiter bei Abkühlung auf den absoluten Nullpunkt — also bei $-273\text{ }^{\circ}\text{C}$ — auf den Wert 0 Ohm (Supraleitung), während der Widerstand eines reinen Halbleiterstoffes dabei unendlich groß wird. Um das zu erklären, soll kurz auf den Aufbau der Halbleiterstoffe eingegangen werden.

In der Halbleitertechnik müssen wir unterscheiden zwischen

a) dem Halbleiterstoff,

d.h. entweder der chemisch reine Grundstoff (z.B. chemisch reines Germanium oder Silizium) oder ein durch bestimmte Beimengungen gewonnener Stoff = dotierter Halbleiterstoff (z.B. P-Germanium oder N-Germanium) und

b) dem Halbleiter-Bauelement,

d.h. ein aus Halbleiterstoffen zusammengesetztes Bauelement mit z.B. Gleichrichtereigenschaft (Diode) oder Verstärkereigenschaft (Transistor).

2.1.1. Kristallaufbau der Halbleiterstoffe

Halbleiterstoffe bestehen aus sogenannten **Kristallgittern**, wobei die Elektronen auf den äußeren Elektronenschalen (Valenzelektronen) der Atome fest im Kristallgitter gebunden sind. Man spricht dabei von einer vollkommenen **Elektronenpaarbindung**: Je ein Elektron zweier benachbarter Atome vereinigen sich zu einem Elektronenpaar (Va-

lenzbrücke). Diese Elektronenpaare haben untereinander eine wesentlich festere Bindung als die Elektronen an den Atomkern; sie können daher auch nur durch große Energiezufuhr getrennt werden. Die zu Elektronenpaaren gebundenen Elektronen sind also nicht frei beweglich und können somit auch nicht zum Ladungstransport herangezogen werden (keine Leitfähigkeit für den elektrischen Strom). In vereinfachter Darstellung soll das am Beispiel des Halbleiterstoffes Germanium erläutert werden.

Die Germanium-Atome enthalten auf der äußeren Elektronenschale 4 Elektronen; Germanium ist somit ein vierwertiges chemisches Element. Diese 4 Valenzelektronen können sich jetzt mit den entsprechenden äußeren Elektronen benachbarter Germanium-Atome zu Elektronenpaaren binden. In diesem Zustand — der beim absoluten Nullpunkt gegeben ist — ist Germanium ein Nichtleiter, da keine freien Elektronen vorhanden sind (vgl. Abb. 2.1).

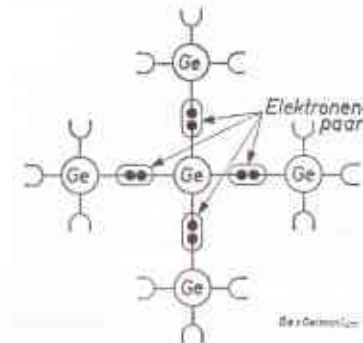


Abb. 2.1 — Kristallgitter des Germaniums bei $-273\text{ }^{\circ}\text{C}$

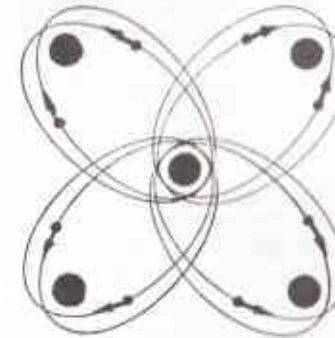


Abb. 2.2 — Elektronenpaarbindung

Die Elektronen der in der Darstellung außenliegenden Germanium-Atome sind natürlich ebenfalls mit Elektronen benachbarter Atome zu Elektronenpaaren vereint. Da Kristallgitter eine räumliche Struktur besitzen (Germanium z.B. Tetraeder), ist die flächige Darstellung hier stark vereinfacht worden.

Man kann sich die Elektronenpaarbindung etwa so vorstellen, wie es nebenstehende Abbildung zeigt. Je zwei Elektronen benachbarter Atome bilden ein Paar, das auf einer gemeinsamen Bahn die beiden zugehörigen Atomkerne umkreist. Daraus ergibt sich eine sehr feste Bindung zwischen den Atomen. Findet eine solche Elektronenpaarbindung zwischen allen Atomen eines Stoffes statt, so ist dieser sehr hart und elektrisch nicht leitfähig.

Jedes Germaniumatom ist innerhalb des Kristallgitters mit je 4 benachbarten Germaniumatomen durch je ein Elektronenpaar fest gebunden. Da sämtliche Valenzelektronen zur Paarbindung verwendet werden, spricht man auch von vollkommener Elektronenpaarbindung. Zu den chemischen Elementen, die zur vollkommenen Elektronenpaarbindung neigen, gehören Kohlenstoff, Silizium und Germanium. Die feste Bindung ist an der großen Härte erkennbar (Kohlenstoff mit vollkommener Elektronenpaarbindung = Diamant = härtester Stoff).

Im Gegensatz zu den Halbleiterstoffen sind bei den reinen Metallen mehr Nachbaratome als Valenzelektronen vorhanden. Hier findet also nicht jedes Valenzelektron einen Partner zur Elektronenpaarbindung (unvollkommene Elektronenpaarbindung). Aus diesem Grunde sind je Metallatom ein oder mehrere Valenzelektronen frei beweglich. Daher leiten Metalle auch bei niedrigen Temperaturen gut — ja, sie neigen sogar dazu, in der Nähe des absoluten Nullpunkts ihren elektrischen Widerstand ganz zu verlieren (Supraleitung). Bei Raumtemperatur entfällt bei Metallen auf jedes Atom etwa ein freies Elektron, bei reinen Halbleiterstoffen dagegen erst auf 10^9 (= 1 Milliarde) Atome ein freies Elektron.

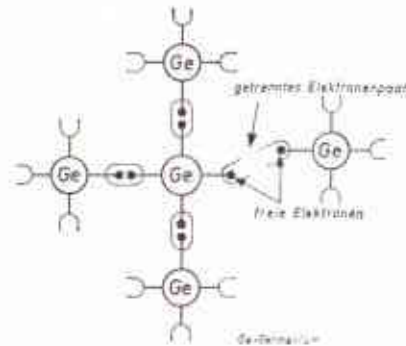


Abb. 2.3 — Kristallgitter bei Erwärmung

Wird das Gefüge jetzt erwärmt, so geraten die Kristallgitter in Schwingungen. Dabei kommt es zur Trennung einzelner Elektronenpaare. Je größer die zugeführte Wärmeenergiemenge ist, desto mehr Elektronenpaare werden getrennt (Abb. 2.3).

Da die nicht zu Elektronenpaaren gebundenen Elektronen frei sind, wird der Stoff nun elektrisch leitend. Man spricht in diesem Zusammenhang von der **Eigenleitfähigkeit** eines Halbleiters.

2.1.2. N-Leiter/P-Leiter (Störstellenleiter)

Das Leitvermögen für den elektrischen Strom schwankt bei reinen Halbleiterstoffen in Abhängigkeit von der Temperatur in weiten Grenzen, nämlich zwischen dem eines Isolators und dem eines metallischen Leiters. Das würde z.B. in der Raumfahrttechnik zu erheblichen Schwierigkeiten führen, da bei der Raumfahrt mit sehr großen Temperaturschwankungen gerechnet werden muß. Es ist jedoch möglich, die Leitfähigkeit von Halbleiterstoffen durch „Verunreinigungen“ zu verbessern. Als „Verunreinigungen“ wählt man che-

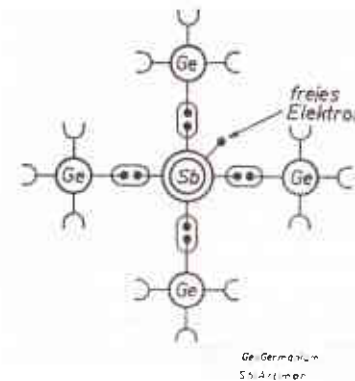


Abb. 2.4 — N-Leiter

mische Elemente, deren Atome auf der äußeren Elektronenschale entweder ein Elektron mehr oder weniger als die Atome des Halbleiterstoffs enthalten¹⁾. Wird z.B. dem Germanium (4 Valenzelektronen) Antimon (5 Valenzelektronen) zugesetzt, so bilden sich Kristalle, bei denen ein Elektron (nämlich das 5. Valenzelektron des Antimons) keinen Partner zur Bildung eines Elektronenpaares findet (vgl. Abb. 2.4).

Dieses eine Elektron ist frei beweglich und kann zum Ladungstransport herangezogen werden. Da hier **negativ geladene Elektronen** zum Ladungstransport dienen, bezeichnet man einen solchen Halbleiterstoff auch als **N-Leiter**. Die Stelle, an der durch den Einbau eines Fremdatoms ein Elektron frei geworden ist, wird als **Störstelle** bezeichnet.

Wird dem Germanium (4 Valenzelektronen) z.B. Indium (3 Valenzelektronen) beigegeben, ergeben sich Kristalle, denen zur vollkommenen Elektronenpaarbindung ein Elektron fehlt. Der Platz, an dem ein Elektron fehlt, wird als

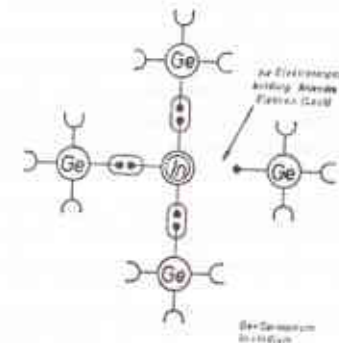


Abb. 2.5 — P-Leiter

„**Loch**“ oder Defektelektron oder Störstelle bezeichnet. Ein „Loch“ hat das Bestreben, ein Elektron einzufangen, um damit wieder ein vollständiges Kristallgitter aufbauen zu können. Es wirkt also etwa wie eine positive Ladung; denn auch positive Ladungen ziehen Elektronen an. Da hier zur vollkommenen Elektronenpaarbindung **Elektronen fehlen**, bezeichnet man diesen Stoff als **P-Leiter**. Das Leitvermögen eines P-Leiters wird

häufig auch als „**Löcherleitung**“ angegeben. Zum besseren Verständnis wollen wir nachfolgend die Elektronen als negative Ladungen (—) und die „Löcher“ als positive Ladungen (+) auffassen. Die Leitfähigkeit eines N- oder P-dotierten Halbleiterstoffs heißt **Störstellenleitung**.

¹⁾ Das Verunreinigen eines reinen Halbleiter-Elements wird in der Technik als **Dotieren** bezeichnet.

Die Frage „Was ist ein Halbleiter?“ kann also wie folgt beantwortet werden:

Die Halbleiterstoffe Germanium und Silizium sind von Natur aus eigentlich Nichtleiter. Ihr Leitvermögen bei der Raumtemperatur 20 °C beruht darauf, daß infolge der Wärmebewegung im Kristallgitter Elektronen frei werden. Durch Beimengung (Dotieren) von 5- oder 3wertigen Elementen kann die elektrische Leitfähigkeit eines Halbleiterstoffes erheblich erhöht werden. Dabei ergibt die Beimengung von 5wertigen Elementen zusätzliche freie Elektronen (N-Leiter), die Beimengung 3wertiger Elemente sogenannte Löcher (P-Leiter). Vereinfacht kann man sagen, daß ein N-Leiter frei bewegliche negative Ladungen (Elektronen) und ein P-Leiter frei bewegliche positive Ladungen (Löcher) enthält.

Die ausschließlich auf Wärme zurückzuführende Leitfähigkeit eines undotierten Halbleiters wird als Eigenleitfähigkeit, die durch Dotieren bewirkte bessere Leitfähigkeit als Störstellenleitung bezeichnet.

2.1.3. Das Temperaturverhalten von Halbleitern

Halbleiter zeigen im Vergleich zu metallischen Leitern eine viel stärkere Abhängigkeit ihres elektrischen Leitvermögens von der Temperatur. Wie schon festgestellt wurde, müßten reine Halbleiterstoffe wie Germanium und Silizium eigentlich Isolatoren sein. Daß sie trotzdem den elektrischen Strom leiten, hängt mit der Wärmebewegung im Kristallgitter der Halbleiter zusammen; denn wir betrachten die Vorgänge ja üblicherweise bei der normalen Raumtemperatur von + 20 °C. Diese Temperatur liegt aber schon sehr weit über dem absoluten Nullpunkt von -273 °C. Beim absoluten Nullpunkt wären reine Halbleiterstoffe tatsächlich Isolatoren.

Vielleicht haben Sie schon erkannt: Wenn ein Stoff sein Leitvermögen bei einer Erwärmung vom absoluten Nullpunkt auf die Raumtemperatur vom Nichtleiter bis zu einer gut meßbaren Leitfähigkeit verändert, dann muß sein Leitvermögen für den elektrischen Strom stärker temperaturabhängig sein als bei üblichen Isolatoren. Wir kennen in der allgemeinen Elektrotechnik den sogenannten Temperaturbeiwert α für die temperaturabhängige Widerstandsänderung eines Leitermaterials. Für metallische

Leiter beträgt α im Mittel $+ 0,004 \frac{1}{^{\circ}\text{C}}$; d.h., bei Erwärmung eines Metalldrahts mit 1 Ohm Widerstand um 1 °C vergrößert sich sein Widerstand um 0,004 Ohm, also um 0,4 %.

Bei Halbleitern nimmt der Widerstand mit steigender Temperatur ab. Ihr Temperaturbeiwert beträgt etwa $\alpha = -0,04 \frac{1}{^{\circ}\text{C}}$, also -4 %/°C. Gegenüber dem Temperaturbeiwert der Metalle ist er also 10mal so groß. Eine Temperaturerhöhung um etwa 9 °C kann bereits zur Verminderung des Widerstands auf den halben Wert führen.

Wir wollen dazu einen sehr anschaulichen Versuch machen. Nach Abb. 2.6 wird mit einem Transformator eine Spannung von etwa 440 Volt erzeugt. Diese Spannung wird über ein Amperemeter (Meßbereich ca. 2 Ampere) an ein Röhrchen aus Laborglas gelegt. Wir wissen, daß Glas zu den Isolatoren zählt; beim Anlegen der Spannung fließt also auch kein Strom. Der Zeiger des Amperemeters schlägt nicht aus. Wird jedoch das Glasröhrchen mit einem Gasbrenner auf einige 100 °C erhitzt, so zeigt sich folgendes: Der Strommesser zeigt plötzlich einen zunehmend stärkeren Strom an. Das Glasröhrchen glüht dabei immer heller auf und wird schließlich so heiß, daß es wie ein Sicherungsdraht durchschmilzt. Die Stromstärke beträgt dabei ca. 1 bis 2 Ampere.

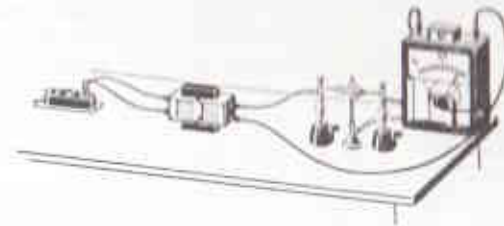


Abb. 2.6 — Die Leitfähigkeit von Glas in Abhängigkeit von der Temperatur

Wie ist das zu erklären? Nun, im Glas geht an sich das gleiche vor wie in einem Halbleiterstoff. Durch die bei Erwärmung zunehmende Bewegung der Atome werden Elektronen freigesetzt und diese freien Elektronen durch die angelegte Spannung bewegt (elektr. Strom). Infolge des Stromflusses wird aber elektrische Energie verbraucht und in Wärme umgewandelt. Das Glasröhrchen erwärmt sich stärker; damit werden zusätzliche Elektronen frei und die Stromstärke wächst weiter.

An diesem Versuch kann man erkennen, daß ein Isolator nur bei normaler Raumtemperatur als Nichtleiter anzusehen ist. Die Zahl der freien Elektronen in Isolatoren ist bei Raumtemperaturen noch so gering, daß wir beim Anlegen einer Spannung mit üblichen Meßinstrumenten keinen Strom nachweisen können. In Halbleiterstoffen reicht aber die Erwärmung vom absoluten Nullpunkt auf die Raumtemperatur 20 °C bereits aus, um eine meßtechnisch nachweisbare Zahl Elektronen freierwerden zu lassen.

Um das Temperaturverhalten von Halbleitern zu untersuchen, hier ein weiterer Versuch. Dazu verwenden wir eine Universaldiode auf Germaniumbasis, z.B. AA 118 oder eine ähnliche Diode, und schließen sie an ein Ohmmeter oder Vielfachinstrument, dessen Widerstandsmessbereich auf x 1k geschaltet ist. Die Diode soll so an das Meßinstrument geschaltet werden, daß sie mit der im Instrument eingebauten Stromquelle in Sperrichtung betrieben wird. Die Sperrichtung ist an der Anzeige eines hohen Widerstandswerts erkennbar. Der Sperrwiderstand einer Germaniumdiode beträgt etwa 1 bis 2 Megohm. Halten wir jetzt die Diode im eingeschalteten Zustand in den warmen

Luftstrom eines Haartrockners oder Heizlüfters (Vorsicht! Nicht zu stark erwärmen, da sonst die Diode zerstört werden kann!), so ist zu erkennen:

Der Widerstandswert nimmt ab und sinkt schnell auf ca. 100 k Ω oder weniger, also auf $\frac{1}{10}$ bis $\frac{1}{20}$ des Widerstands bei normaler Umgebungstemperatur. Lassen wir die Diode wieder abkühlen (ggf. kalt anblasen), so steigt der Widerstand langsam wieder auf den ursprünglich hohen Wert.

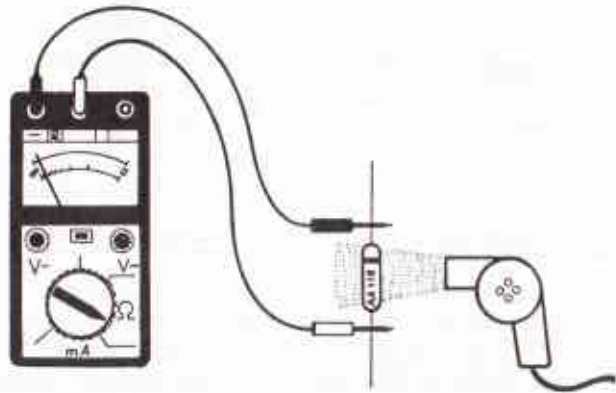


Abb. 2.7. — Nachweis des Temperaturverhaltens einer Diode

Eine Erwärmung von Halbleiter-Bauelementen muß aber nicht nur von außen durch eine hohe Umgebungstemperatur bedingt sein. Sie kann vielmehr auch durch die im Betrieb erzeugte Verlustwärme im Bauelement selbst hervorgerufen werden. Stellen wir uns doch einmal vor, ein Halbleiter-Bauelement — z.B. eine Diode — wird an einer festen Gleichspannung von 10 V betrieben; sie besitzt bei Raumtemperatur einen bestimmten, meßbaren Widerstand, z.B. 100 Ohm. Über unsere Diode fließt jetzt ein Strom von

$$I = \frac{U}{R} = \frac{10 \text{ V}}{100 \Omega} = 0,1 \text{ A.}$$

Am Diodenwiderstand wird aufgrund der Verlustleistung von

$$P = I^2 \cdot R = 0,1^2 \cdot 100 = 1 \text{ Watt}$$

Wärme erzeugt. Diese Wärme setzt aber — wenn sie nicht schnell genug abgeführt wird — den Widerstand des Halbleitermaterials

herab. Bei halbiertem Widerstand, also 50 Ohm, wird der Strom auf den doppelten Wert, also 0,2 A, ansteigen. Dadurch steigt die Verlustleistung auf

$$P_t = I_t^2 \cdot R_t = 0,2^2 \cdot 50 = 2 \text{ Watt!}$$

Großere Verlustleistung bedeutet aber stärkere Erwärmung; sie setzt den Widerstand des Halbleitermaterials herab. Bei weiter verringertem Widerstand fließt ein noch stärkerer Strom... nun, Sie ahnen vielleicht schon, wie's weitergeht? Unser Halbleiter-Bauelement erwärmt sich innerhalb kurzer Zeit so stark, daß es zerstört wird. Aus diesem Grund müssen wir bei allen Versuchen und Schaltungen mit Halbleiter-Bauelementen dafür sorgen, daß ein unzulässig hoher Strom vermieden wird. Wie können wir das tun? Kommen wir noch einmal auf unser Zahlenbeispiel zurück:

Unser Halbleiter-Bauelement hat bei Raumtemperatur einen Widerstand von 100 Ohm und liegt an einer Spannung von 10 Volt. Jetzt schalten wir einmal einen gleich großen Widerstand — also 100 Ohm — in Reihe zum Halbleiter-Bauelement. Damit für das Bauelement gleiche Betriebsbedingungen herrschen, muß die Spannung auf 20 Volt erhöht werden.

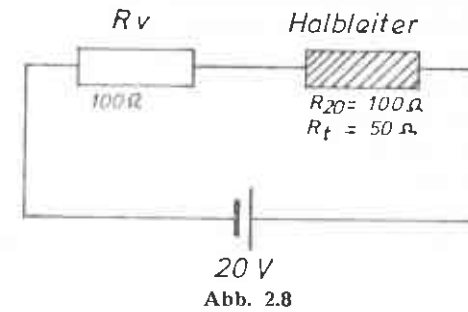


Abb. 2.8

Nach dem Ohmschen Gesetz fließt auch in dieser Schaltung zunächst ein Strom von

$$I = \frac{U}{R_{20} + R_v} = \frac{20 \text{ V}}{200 \Omega} = 0,1 \text{ A.}$$

Verringert sich der Halbleiterwiderstand durch Erwärmung auf $R_t = 50 \text{ Ohm}$, so beträgt die Stromstärke

$$I_t = \frac{U}{R_t + R_v} = \frac{20 \text{ V}}{150 \Omega} = 0,133 \text{ A.}$$

Die Verlustleistung am Halbleiter-Bauelement ergibt sich zu

$$P_L = I_L^2 \cdot R_L = 0,133^2 \cdot 50 = 0,0177 \cdot 50 = 0,89 \text{ Watt}$$

Die Verlustleistung ist also anstatt größer, wie im ersten Beispiel, jetzt sogar kleiner geworden. Somit kann sich unser Halbleiter-Bauelement nicht mehr so stark erwärmen. Man kann sagen: Der Transistor ist bei 20 °C leistungsgangepaßt.

Ein Halbleiter-Bauelement kann nicht nur durch eine hohe Umgebungstemperatur gefährdet werden, es erwärmt sich auch im Betrieb durch Stromwärme selbst. Diese Eigenerwärmung kann man außer durch schnelle Wärmeableitung mit Hilfe von Kühlflächen auch durch geeignet bemessene Vorwiderstände im Betriebsstromkreis herabsetzen.

Soll auf eine genaue Berechnung des Widerstandswerts verzichtet werden, so wird der Vorwiderstand genauso groß wie der Widerstand des Halbleiter-Bauelements bei Raumtemperatur gewählt.

Diese Erkenntnis ist für die Herstellung von Halbleiter-Schaltungen von großer Wichtigkeit. Wird dem Temperaturverhalten der Halbleiter-Bauelemente zu wenig Beachtung geschenkt, so ist die richtige Funktion einer Schaltung in Frage gestellt. U.U. werden Bauteile der Schaltung im Betrieb sogar schnell zerstört. Auch für unsere künftigen Schaltungen ist diese Tatsache zu berücksichtigen.

Da vom Hersteller von Halbleiter-Bauelementen die Höchstwerte für Spannung, Strom und Verlustleistung angegeben werden, sind die Widerstandswerte der erforderlichen Vorwiderstände mit Hilfe der Gesetzmäßigkeiten für den Gleichstromkreis leicht zu errechnen. Am einfachsten wird aber die richtige Auslegung einer Schaltung mit Hilfe der Kennlinien für die verwendeten Bauteile vorgenommen. Damit wir im Umgang mit Kennlinien sicherer werden, soll noch einmal das Wichtigste dazu wiederholt werden.

2.1.4. Die Berücksichtigung der zulässigen Verlustleistung eines Bauelements im Kennlinienfeld

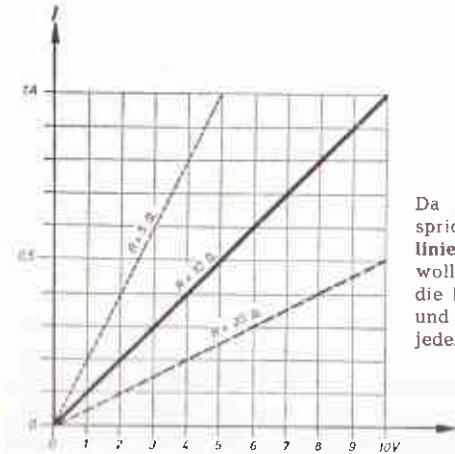
Die für Bauteile üblichen Kennlinien stellen normalerweise die Abhängigkeit einer Größe (z.B. des Stroms) von einer anderen Größe (z.B. der Spannung) dar. Am gebräuchlichsten ist die Spannungs-Strom-Kennlinie. Sie veranschaulicht, wie sich der durch das betreffende Bauteil fließende Strom verhält, wenn man die anliegende Spannung ändert. Eine solche Kennlinie zeigt damit das Widerstandsverhalten des Bauelements. Sehr einfach ist die Spannungs-Strom-Kennlinie für einen ohmschen Widerstand. Steigern wir an einem Widerstand von 10 Ohm die angelegte Spannung von 0 bis 10 Volt, so steigt der Strom von 0 bis 1 Ampere. Die einzelnen Zwischenwerte

berechnen wir nach dem Ohmschen Gesetz $I = \frac{U}{R}$ und stellen sie in Tabellenform zusammen:

$R = 10 \text{ Ohm}$

U	0	1	2	3	4	5	6	7	8	9	10	V
I	0	0,1	0,2	0,3	0,4	0,5	0,6	0,7	0,8	0,9	1,0	A

Tragen wir die Werte jetzt in ein Diagramm, so ergibt sich folgende Kennlinie (ausgezogene Linie!).



Da die Kennlinie geradlinig verläuft, spricht man auch von der linearen Kennlinie eines ohmschen Widerstands. Wir wollen in das gleiche Diagramm ebenso die Kennlinien für die Widerstände 5 Ohm und 20 Ohm eintragen und stellen dazu für jeden Widerstand eine Wertetabelle auf:

Abb. 2.9 — Widerstandskennlinien

$R = 5 \text{ Ohm}$

U	0	1	2	3	4	5	6	7	8	9	10	V
I	0	0,2	0,4	0,6	0,8	1,0	1,2	1,4	1,6	1,8	2,0	A

$R = 20 \text{ Ohm}$

U	0	1	2	3	4	5	6	7	8	9	10	V
I	0	0,05	0,1	0,15	0,2	0,25	0,3	0,35	0,4	0,45	0,5	A

Aus dem Kennlinienverlauf der drei Widerstände lassen sich folgende Rückschlüsse ziehen: Die Spannungs-Strom-Kennlinie eines ohmschen Widerstands ist linear. Aus der Steigung der Kennlinie kann man auf die Größe des Widerstands schließen. Für einen großen Widerstand ergibt sich eine geringe Steigung (flacher Verlauf), für einen kleinen Widerstand eine große Steigung (steiler Verlauf) der Kennlinie.

Bei den Meßübungen mit Oszilloskopen sind bereits Widerstandskennlinien für verschieden große Widerstände dargestellt worden. Die dort ermittelten Diagramme werden hier durch Berechnung bestätigt.

Wie können wir nun die zulässige Verlustleistung und damit die Erwärmung des Bauelements im Spannungs-Strom-Diagramm berücksichtigen? Die elektrische Leistung ergibt sich aus dem Produkt Spannung mal Strom. Als Formel geschrieben: $P = U \cdot I$. In einem Spannungs-Strom-Diagramm kann man eine Linie darstellen, die eine bestimmte Leistung — in unserem Fall die zulässige Verlustleistung — kennzeichnet. Wenn nämlich eine bestimmte Leistung gegeben ist,

dann gehört zu einem bestimmten Spannungswert ein ganz bestimmter Stromwert. Anders ausgedrückt: Je größer die Spannung wird, desto kleiner muß der Strom werden, damit die Leistung gleichbleibt. Die

zugehörigen Werte können mit Hilfe der Leistungsformel $I = \frac{P}{U}$ berechnet werden.

Zur Übung verwenden wir unser Diagramm mit den Kennlinien der drei Widerstände 5, 10 und 20 Ohm und nehmen an, daß die Widerstände mit höchstens 1 Watt belastbar sind. Für 1 Watt und Spannungswerte zwischen 0 und 10 Volt ergibt sich folgende Wertetabelle:

U	0	1	2	3	4	5	6	7	8	9	10	V
I	∞	1	0,5	0,33	0,25	0,2	0,167	0,143	0,125	0,111	0,1	A

Tragen wir die zugehörigen Werte der Tabelle in ein Spannungs-Strom-Diagramm und verbinden die Punkte untereinander, so ergibt sich eine Linie in Form einer Hyperbel (d.h. eine Kurve, die zwischen zwei sich schneidenden Geraden — in unserem Fall die Spannungsachse und die Stromachse — verläuft und sich diesen Geraden im Unendlichen nähert); man spricht daher auch von der **Leistungshyperbel**. Die Leistungshyperbel im sogenannten Kennlinienfeld macht uns die Einhaltung der zulässigen Verlustleistung sehr einfach. Für alle Punkte des Kennlinienfeldes, die unterhalb der Leistungshyperbel liegen, ist die Leistung kleiner als die zulässige Leistung. Dagegen ist die Leistung für alle Punkte oberhalb der Leistungshyperbel zu groß.

Für unser Diagramm ergibt sich das folgende Bild:

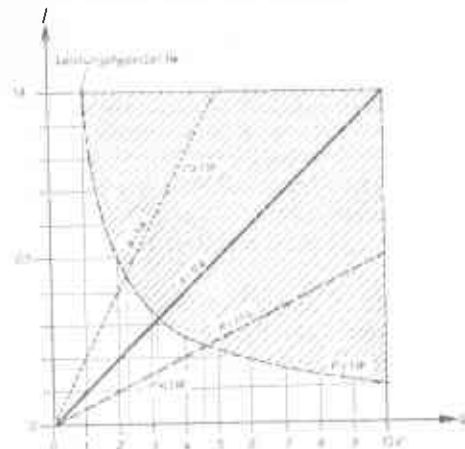


Abb. 2.10 — Leistungshyperbel im Widerstands-Kennlinienfeld

Aus dem Widerstands-Kennlinienfeld mit Leistungshyperbel sind jetzt sehr schnell und einfach die Grenzwerte abzulesen. Damit die an den Widerständen wirksame Leistung nicht größer als 1 Watt wird, darf die anliegende Spannung für den

5-Ohm-Widerstand 2,23 Volt,

10-Ohm-Widerstand 3,16 Volt und den

20-Ohm-Widerstand 4,47 Volt

nicht überschreiten.

Genauso wie für die Widerstände wird mit den Kennlinien der Halbleiterbauelemente verfahren und in das jeweilige Kennlinienfeld auch die Leistungshyperbel eingetragen und damit sichergestellt, daß die für das Bauelement zulässige Verlustleistung nicht überschritten wird.

2.2. Einfache Halbleiter-Bauelemente

2.2.1. NTC-Widerstände (Heißleiter/Thermistoren)

NTC-Widerstände verringern ihren Widerstand mit zunehmender Temperatur; ihr Widerstands-Temperaturbeiwert α ist also negativ. Die Bezeichnung NTC-Widerstand ist aus dem Begriff negativer Temperatur-Coeffizient abgeleitet. Der NTC-Widerstand besteht also aus einem Stoff, der im heißen Zustand besser leitet als im kalten; er führt aus diesem Grund häufig auch die Bezeichnung Heißleiter oder Thermistor. Man könnte nun einfach Germanium oder Silizium als Widerstandsmaterial für Heißleiter verwenden, da sie ja als Halbleiterstoffe einen negativen Temperaturbeiwert, also Heißleitereigenschaften, besitzen. Dem stehen aber zwei Umstände entgegen: Erstens sind Halbleiterstoffe sehr teuer (1 g reines Germanium kostet etwa 6,— DM und ist fast so teuer wie Gold!), und zweitens können sie wegen ihrer großen Härte und Sprödigkeit nur sehr schwierig zu größeren Bauelementen verarbeitet werden. Daher finden zur Herstellung von Heißleitern Stoffe Verwendung, die billiger sind und sich leichter verarbeiten lassen; hierzu gehören Magnesiumoxid, Titanoxid, oder Magnesium-Nickel-Oxid; sie werden meistens zur Formgebung gesintert.

Heißleiter werden in verschiedensten Bauformen hergestellt, und zwar in Stabform (ähnlich wie Kohleschichtwiderstände), Scheiben- oder Plattenform. Bei einigen Typen ist der Heißleiter auch in ein Glasröhrchen oder einen Glaskolben eingebaut. Besondere Regel-

Heißleiter sind mit einer zusätzlichen Heizwicklung versehen, die im Gehäuse untergebracht ist; Abb. 2.11 zeigt eine Auswahl von Heißleiter-Bauformen (etwa $\frac{1}{2}$ natürl. Größe).

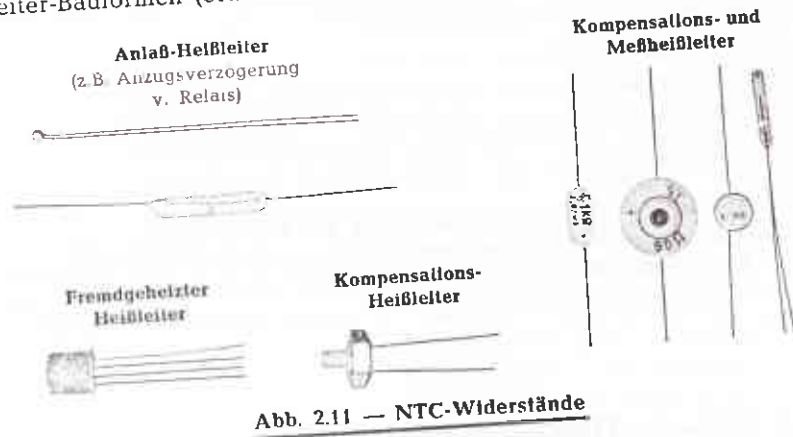


Abb. 2.11 — NTC-Widerstände

Als Schaltzeichen für NTC-Widerstände wird nebenstehendes Symbol verwendet, wobei das Minuszeichen neben dem Formelzeichen der Temperatur (ϑ) auf den negativen Temperaturbeiwert hinweist. In Industrieschaltungen findet man jedoch anstelle ϑ die Angabe NTC neben dem Widerstandssymbol.

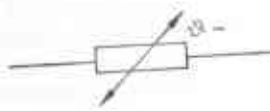


Abb. 2.12

2.2.1.1. Eigenschaften der NTC-Widerstände

Legt man an einen NTC-Widerstand eine Spannung, so fließt wegen seines hohen Kaltwiderstands zunächst ein geringer Strom. Infolge des Verbrauchs elektrischer Energie (Verlustleistung) erwärmt sich der NTC-Widerstand und verringert seinen Widerstandswert. Da er grundsätzlich mit einem Vorwiderstand betrieben werden muß, sinkt die Teilspannung am NTC-Widerstand. Nach einiger Zeit stellt sich ein Gleichgewichtszustand zwischen der durch die Verlustleistung bewirkten Eigenerwärmung des Heißleiters und der an die Umgebung abgegebenen Wärme ein; man nennt diesen Zustand „stationär“. Wird die Spannung weiter erhöht, so folgt die Erwärmung des Heißleiters infolge seiner Wärmeträgheit nur langsam nach, bis sich wiederum ein stationärer Zustand einstellt. Das Widerstandsverhalten eines NTC-Widerstands ist nicht linear. Wir können daher bei der

Behandlung der Grundlagen nicht auf Näherungsformeln für die Berechnung des Widerstandsverhaltens eingehen. Ein Spannungs-Strom-Diagramm veranschaulicht aber das Widerstandsverhalten.

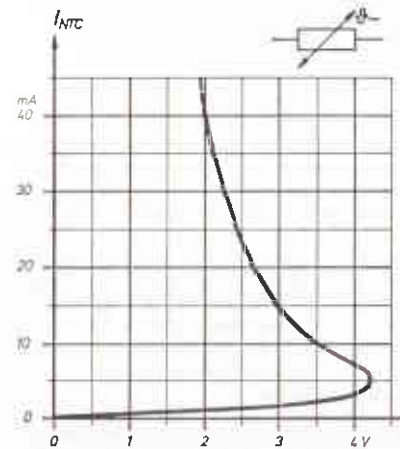


Abb. 2.13 — Kennlinie eines NTC-Widerstands

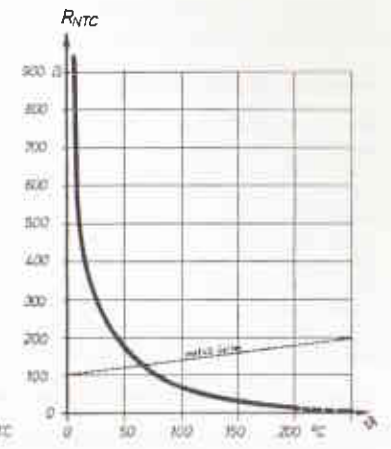


Abb. 2.14 — Temperaturabhängigkeit des NTC-Widerstands

Die Kennlinie läßt erkennen, wie mit steigender Spannung zunächst nur ein geringer Stromanstieg bewirkt wird. Sobald sich aber der Heißleiter erwärmt hat, sinkt sein Widerstand. Damit fällt die an ihm liegende Spannung, während der Strom ansteigt. Der zulässige Temperaturbereich geht bei den handelsüblichen NTC-Widerständen bis etwa $+200^{\circ}\text{C}$. Im Mittel verändert sich dabei der Widerstandswert im Verhältnis $100 : 1$; d.h., er beträgt bei $+200^{\circ}\text{C}$ nur noch $1/100$ des Widerstands bei 20°C . Je größer der Kaltwiderstand R_{20} eines NTC-Widerstands ist, desto größer ist das Verhältnis von Kalt- zu Warmwiderstand. Als wichtigste Kennwerte für NTC-Widerstände gelten

- a) der Kaltwiderstand bei 20°C (R_{20}),
- b) der Warmwiderstand bei der zulässigen Höchsttemperatur,
- c) der Widerstands-Temperaturbeiwert α (er beträgt für NTC-Widerstände zwischen $-0,03$ und $-0,055$ $1/^{\circ}\text{C}$),
- d) die zulässige Verlustleistung und
- e) die zulässige Höchsttemperatur.

2.2.1.2. Anwendungsbeispiele für NTC-Widerstände

Da sich das Temperatur-Widerstandsverhalten eines Heißleiters entgegengesetzt zu dem der metallischen Leiter verhält, dient er vorwiegend der Temperaturkompensation in Stromkreisen mit metallischen Leitern; dazu zwei Beispiele:

Beispiel 1: Bildhöhenstabilisierung im Fernsehgerät

Die zur Ablenkung des Elektronenstrahls in einer Fernseh-Bildröhre dienenden Spulen (magnetische Ablenkung) bestehen aus vielen Windungen Kupferdraht. Infolge der Erwärmung des Fernsehgeräts im Betriebszustand nimmt der Widerstand der Kupfer-Erwicklung zu, damit wird der Strom kleiner. Ein geringerer Strom erzeugt aber eine weniger kräftige Magnetfeld. Der Elektronenstrahl wird also nicht mehr so weit abgelenkt wie im kalten Zustand des Geräts. Das äußert sich sehr kraß durch ein Zusammenschrumpfen des Fernsehbilds — vor allem in der Bildhöhe. Schaltet man aber in Reihe zu den Ablenkspulen einen NTC-Widerstand, der auf die Erwärmung des Fernsehgeräts abgestimmt ist, so verringert dieser seinen Widerstand um den gleichen Betrag, wie die Kupferwicklungen ihren Widerstand erhöhen. Die Bildhöhe bleibt damit unbeeinträchtigt von der Betriebstemperatur des Geräts.

Beispiel 2: Temperaturkompensation bei Meßinstrumenten

Die Spulen von Meßinstrumenten (z.B. Zeigerinstrument mit eingebautem Drehspulmeßwerk) werden aus dünnem Kupferdraht gewickelt; der Temperaturbeiwert für Kupfer ist positiv. Wird das Meßinstrument bei unterschiedlichen Umgebungstemperaturen verwendet, so ändert sich der Innenwiderstand des Meßwerks und damit die Anzeigegenauigkeit. Eine Temperaturänderung um 3°C hat bereits eine Widerstandsänderung von mehr als 1% bei der Kupferdrahtspule zur Folge. Vor allem in Meßinstrumente der Güteklassen 0,1 bis 0,5 baut man daher zur Temperaturkompensation NTC-Widerstände ein und gleicht damit die Widerstandsänderungen der Kupferwicklung aus. Da aber NTC-Widerstände viel stärker temperaturabhängig sind als metallische Leiter, muß man die Widerstandsänderung der NTC-Widerstände durch Parallelschalten von Nebenwiderständen eingrenzen!

Die für die vorgenannten Zwecke verwendeten NTC-Widerstände werden als **Kompensations-Heißleiter** bezeichnet; sie dienen auch zur Stabilisierung von Transistorschaltungen.

Als weitere Anwendungsmöglichkeit bieten sich Heißleiter zur Unterdrückung von Einschaltstromstößen oder zur Anzugs- bzw. Abfallverzögerung von Relais an. Speziell für diese Zwecke hergestellte NTC-Widerstände werden als **Anlaß-Heißleiter** bezeichnet. Auch dazu zwei Beispiele:

Beispiel 1: Unterdrückung von Einschaltstromstößen

Schaltet man Stromkreise mit Glühlampen, Elektronenröhren oder Motoren, so ist der Einschaltstromstoß im Verhältnis zur Betriebsstromstärke sehr groß. Dabei können vorgeschaltete, auf die Nennstromstärke abgestimmte Sicherungen auslösen oder Netzsprungspannungsschwankungen auftreten. Schaltet man in diese Stromkreise NTC-Widerstände mit kurzer Erwärmungszeit, so wird der Einschaltstromstoß wegen ihres hohen Kaltwiderstands stark herabgesetzt. Soll der Heißleiter nicht im Stromkreis

belassen werden, so kann man ihn nach dem Einschalten des Stromkreises überbrücken. Glühlampen hoher Leistungsaufnahme werden teilweise mit eingebautem Heißleiter geliefert.

Beispiel 2: Anzugsverzögerung von Relais

Größere Verzögerungen der Anzugs- oder Abfallzeit eines Relais lassen sich mit den üblichen Mitteln, wie Kurzschlußwicklungen oder Kondensatoren, nur schwierig und aufwendig verwirklichen. Man geht daher in zunehmendem Maße dazu über, eine Verzögerung mit einem NTC-Widerstand zu bewirken. Wird der Relaisstromkreis in nebenstehender Schaltung geschlossen, so erhält das Relais wegen des hohen Kaltwiderstands zunächst Fehlstrom. Erst nachdem sich der Heißleiter erwärmt hat, steigt der Strom; das Relais zieht an. Der Heißleiter wird meistens durch einen Arbeitskontakt des Relais überbrückt, damit er sich bis zum nächsten Einschaltvorgang abkühlen kann.

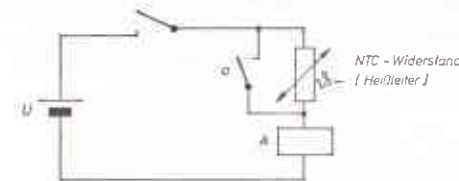


Abb. 2.15 — Anzugsverzögerung eines Relais durch einen NTC-Widerstand

Das Verhalten eines NTC-Widerstands als Anlaß-Heißleiter wird durch folgenden einfachen Versuch veranschaulicht: In einem Stromkreis werden ein NTC-Widerstand und eine Glühlampe hintereinandergeschaltet. Der Kaltwiderstand R_{20} des Heißleiters sollte etwa gleich dem Widerstand der Glühlampe im Betriebszustand sein. Der Heißleiter kann durch einen Schalter überbrückt werden.

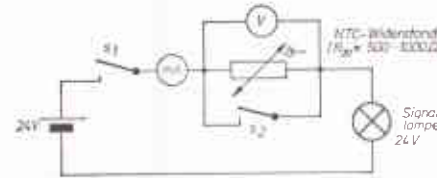


Abb. 2.16 — Einschaltverzögerung einer Signallampe durch einen NTC-Widerstand

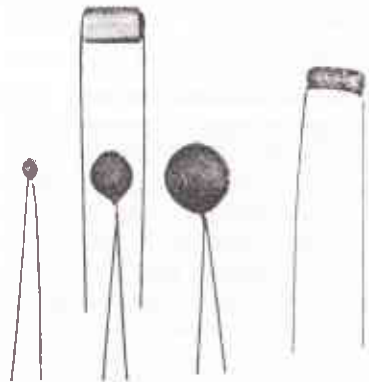
Wird bei geschlossenem Schalter S_2 der Schalter S_1 geschlossen, so leuchtet die Lampe schnell auf. Wird dagegen der Stromkreis bei geöffnetem Schalter S_2 geschlossen, so dauert es einige Zeit, bis die Lampe aufleuchtet. Am eingeschalteten Strommesser kann man das Ansteigen des Stroms, am Voltmeter das Zurückgehen des Spannungsabfalls am NTC-Widerstand verfolgen.

2.2.2. PTC-Widerstände (Kaltleiter)

Grundsätzlich zählen die aus reinen Metallen bestehenden Leiter, wie Kupfer, Aluminium, Silber, Wolfram usw., zu den PTC-Widerständen. Sie leiten bekanntlich im kalten Zustand besser, da ihr Widerstand mit zunehmender Temperatur wächst. Der Temperatureinfluß auf das

Widerstandsverhalten der rein metallischen Leiter ist jedoch verhältnismäßig gering und beträgt im Mittel $+0,004 \text{ } 1/^{\circ}\text{C}$, also $0,4 \text{ } \%$ je $^{\circ}\text{C}$ Temperaturänderung. Man hat durch Untersuchungen festgestellt, daß Stoffe, wie z.B. Bariumtitanat, eine erheblich stärkere temperaturabhängige Widerstandsänderung aufweisen. Der Temperaturbeiwert beträgt hier etwa bis zu $+0,1 \text{ } 1/^{\circ}\text{C}$, also bis zu $10 \text{ } \%$ je $^{\circ}\text{C}$ Temperaturänderung. Früher wurden PTC-Widerstände in Form von Glühlampen mit Metallwendel, heute fast ausschließlich auf Bariumtitanat-Basis in Stab-, Scheiben- oder Kugelform hergestellt. Die Bezeichnung PTC ist aus dem Begriff **positiver Temperatur-Coeffizient** abgeleitet; Abb. 2.17 zeigt einige Bauformen für PTC-Widerstände.

Temperaturmessung und -überwachung



Temperaturfühler



Abb. 2.17 — Bauformen von PTC-Widerständen

Das Schaltzeichen für PTC-Widerstände ist in Abb. 2.18 abgebildet; es unterscheidet sich von dem eines Heißleiters lediglich durch das Pluszeichen hinter dem Formelzeichen der Temperatur, was auf den positiven Temperaturbeiwert hinweist. In Industrieschaltungen findet man anstelle des Formelzeichens $\vartheta +$ auch vielfach die Angabe PTC.

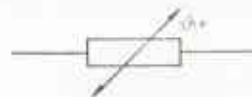


Abb. 2.18

2.2.2.1. Eigenschaften der PTC-Widerstände

Der Kaltwiderstand handelsüblicher PTC-Widerstände beträgt bei 20°C etwa 10 bis 100 Ohm. Wird ein PTC-Widerstand von Strom durchflossen, so erwärmt er sich infolge der Verlustleistung ($P = I^2 \cdot R$). Dabei nimmt der Widerstand zunächst weniger, mit steigender Erwärmung jedoch erheblich schneller zu und steigt bis zur zulässigen Grenztemperatur auf etwa den 100- bis 1000fachen Wert des Kaltwiderstands an. Die zulässige Grenztemperatur beträgt zwischen 150 und 200°C .

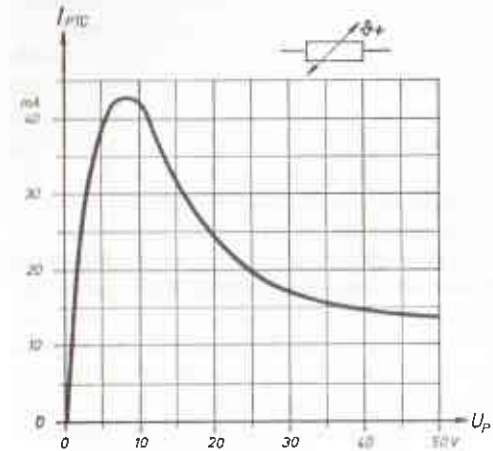


Abb. 2.19 — Kennlinie eines PTC-Widerstands

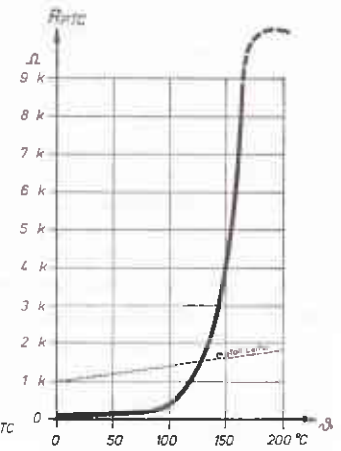


Abb. 2.20 — Temperaturabhängigkeit des PTC-Widerstands

Die Kennlinie und das Temperatur-Widerstands-Diagramm lassen dieses Verhalten deutlich erkennen. Da bei einer kleinen angelegten Spannung nur eine geringe Verlustwärme erzeugt wird, bleibt der Widerstand des Kaltleiters zunächst gering (steiler Verlauf der Kennlinie). Mit steigender Spannung wird aber der PTC-Widerstand wärmer und vergrößert seinen Widerstand; damit fällt jedoch die Stromstärke. Auch bei PTC-Widerständen stellt sich jeweils ein Gleichgewichtszustand zwischen erzeugter Verlustwärme und an die Umgebung abgegebener Wärme ein (stationärer Zustand).

Da PTC-Widerstände teurer als NTC-Widerstände sind und außerdem größeren Fertigungstoleranzen unterliegen, werden für elektronische Temperaturregelaufgaben bevorzugt NTC-Widerstände eingesetzt.

2.2.2.2. Anwendung von PTC-Widerständen

PTC-Widerstände werden überwiegend als sogenannte **Temperaturfühler** verwendet. Da sie sehr klein herstellbar sind, bereitet ihre Unterbringung keine Schwierigkeiten. Sie sprechen nicht nur auf die Eigenerwärmung (Verlustleistung), sondern ebenso auf eine Fremderwärmung (Umgebungstemperatur) an und können damit als elektronisches Bauelement Meß-, Steuer- und Regelaufgaben übernehmen.

Speziell in der Fernmeldetechnik finden Kaltleiter gewissermaßen als Sicherung oder **Überlastungsschutz** Verwendung. Soll z.B. ein empfindliches Gerät oder Bauelement vor Überlastung geschützt werden, so schaltet man in Reihe dazu einen Kaltleiter. Steigt jetzt aus irgendeinem Grund der Strom, so vergrößert der Kaltleiter infolge der größeren Verlustwärme seinen Widerstand und setzt damit die Stromstärke im Stromkreis wieder herab. Bisher wurden die Ruf- und Signalmaschinen durch Glühlampen vor Überlastung geschützt. Die Lampen übernahmen die Funktion eines PTC-Widerstands. Voraussetzlich werden Ruf- und Signalmaschinen schon bald durch Transistorwandler ersetzt, deren Überlastungsschutz PTC-Widerstände übernehmen könnten. Die Schutzfunktion eines Kaltleiters wird durch folgenden Versuch deutlich:

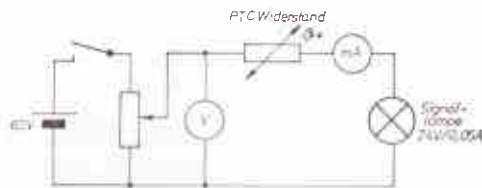


Abb. 2.21 — Überlastungsschutz mit Hilfe eines PTC-Widerstands

Spannungserhöhung wird ihre Helligkeit geringer, u.U. erlischt sie sogar. Trotz einer für die Lampe zu hohen Spannung wird sie nicht überlastet. Der PTC-Widerstand hat seinen Widerstand stark erhöht und damit den Strom im Lampenstromkreis herabgesetzt. Die größte Teilspannung entfällt in diesem Zustand auf den PTC-Widerstand.

2.2.3. Hall- oder Feldplatten (Hallsonden)

Dünne Plättchen aus Halbleiterstoffen zeigen die Eigenschaft, ihren Widerstand unter dem Einfluß eines Magnetfeldes zu vergrößern; es

sind also magnetisch steuerbare Widerstände. Diese Plättchen sind maximal 0,1 mm dick und zum Schutz gegen mechanische Beschädigungen meistens mit einem Kunststoffmantel umgeben. Als Halbleiterstoffe werden hauptsächlich Indiumarsenid, Indiumantimonid oder Indium-Arsen-Phosphor-Verbindungen verwendet. Hall- oder Feldplatten sind so klein herstellbar, daß sie ohne Schwierigkeiten in den schmalen Luftspalt eines magnetischen Kreises eingeführt werden können; Abb. 2.22 zeigt einige Bauformen (ca. $\frac{1}{2}$ natürl. Größe).

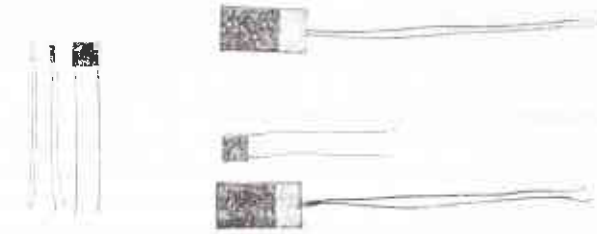


Abb. 2.22 — Hallsonden (magnetisch steuerbare Widerstände)

2.2.3.1. Eigenschaften von Hallsonden

Die Wirkungsweise der Hallsonden beruht auf der Tatsache, daß ein Magnetfeld bewegte Elektronen aus ihrer Bewegungsrichtung ablenkt; denn sobald sich Elektronen bewegen, erzeugen sie selbst ein Magnetfeld (Nachweis des Magnetfeldes um einen stromdurchflossenen Leiter). Zwischen zwei Magnetfeldern besteht aber eine Kraftwirkung — Anziehung oder Abstoßung. Die Richtung, in die ein Strom durch ein äußeres Magnetfeld abgelenkt wird, hängt von der Stromrichtung und der Richtung des Magnetfeldes ab. Die Ablenkrichtung läßt sich mit Hilfe der Linken-Hand-Regel (Motorregel) bestimmen. Denkt man sich den Elektronenstrom in einer dünnen Platte in eine Vielzahl von parallelen Bahnen zerlegt, so werden die Elektronen unter dem Einfluß eines äußeren Magnetfeldes aus ihrer ursprünglichen Bewegungsrichtung abgelenkt. Das wirkt sich in dem sehr dünnen Plättchen so aus, als ob der stromführende Querschnitt kleiner ge-

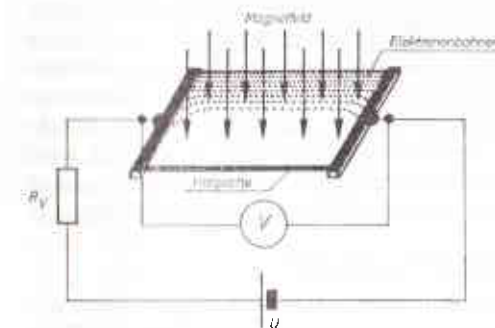


Abb. 2.23 — Wirkungsweise einer Hallplatte (bzw. -sonde)

zahl von parallelen Bahnen zerlegt, so werden die Elektronen unter dem Einfluß eines äußeren Magnetfeldes aus ihrer ursprünglichen Bewegungsrichtung abgelenkt. Das wirkt sich in dem sehr dünnen Plättchen so aus, als ob der stromführende Querschnitt kleiner ge-

worden wäre. Der Widerstand eines Leiters hängt aber wesentlich von seinem Querschnitt ab: Je kleiner der Querschnitt, desto größer der Widerstand. Der Einfluß des äußeren Magnetfeldes auf den Widerstand ist am größten, wenn es senkrecht zur Feldsondenfläche steht. Er steigt zudem fast linear mit der magnetischen Flußdichte B .

Der Einfachheit halber **benutzt man den Spannungsabfall an der Feldplatte als Meßgröße zur Bestimmung der magnetischen Flußdichte B** und kann also folgern: Je stärker das äußere Magnetfeld, desto höher wird bei konstanter Versorgungsspannung der Spannungsabfall an der Hallplatte. Die Hallplatte muß grundsätzlich mit einem Vorwiderstand betrieben werden.

Als Schaltzeichen für eine Hall- oder Feldplatte gilt nebenstehendes Symbol, wobei das Formelzeichen B (magnet. Flußdichte) den Hinweis auf die Beeinflussbarkeit des Widerstands durch eine magnetische Größe gibt.

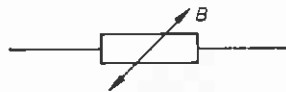


Abb. 2.24

2.2.3.2. Anwendung von Hall- oder Feldplatten

Hallplatten dienen in erster Linie zur Regelung und Steuerung, wobei als Regel- oder Steuergröße ein Magnetfeld dienen muß; man kann Hallplatten auch zur Strommessung verwenden. Der Vorteil dieses Strommeßverfahrens liegt in der galvanischen Trennung zwischen Meßstromkreis und Instrumentenstromkreis. Zur Strommessung mit Hallsonden läßt man den zu messenden Strom durch eine Spule fließen und mißt mit einer Hallsonde indirekt die Stärke des Spulenmagnetfeldes. Bei einer eisenlosen Spule steigt die magnetische Flußdichte linear mit der Stromstärke. Aber auch zur Messung der magnetischen Flußdichte selbst leisten die Hallplatten gute Dienste. Sie lassen die Aufnahme von Magnetisierungskurven für ferromagnetische Werkstoffe in einfacher Weise zu: Der zu untersuchende Werkstoff wird durch eine stromdurchflossene Spule zunehmend magnetisiert und in einem feinen Luftspalt mit einer Hallsonde die magnetische Flußdichte B gemessen.

2.2.3.3. Hallgeneratoren

An dieser Stelle sei noch kurz auf eine besondere Bauform der Hallplatten hingewiesen, die sogenannten Hallgeneratoren. Bei einem Hallgenerator sind an den Flanken einer Hallplatte zusätzlich zwei Elektroden angebracht. Wird der Elektronenstrom in der Platte unter dem Einfluß eines äußeren Magnetfeldes zur Seite abgedrängt, so tritt an dieser Seite eine Elektronenanhäufung, an der anderen Seite ein Elektronenmangel auf; das hat eine elektrische Spannung zur Folge.

Die Höhe dieser Hallspannung hängt vom Umfang der Elektronenverschiebung ab und kann mit einem Spannungsmesser nachgewiesen bzw. für Steuerungs- und Regelungszwecke ausgenutzt werden. Da die Höhe der Hallspannung nahezu ausschließlich von der magnetischen Flußdichte abhängt und die Betriebsspannung nur einen unwesentlichen Einfluß hat, werden für elektronische Steuerungen Hallgeneratoren bevorzugt eingesetzt.

2.2.4. Fotowiderstände

In Halbleiterstoffen können Elektronen durch Energiezufuhr aus ihren Bindungen (Elektronenpaarbindungen) gerissen werden. Am bekanntesten ist die Erhöhung der Zahl an freien Elektronen durch Zufuhr von Wärmeenergie. Da aber **Licht** ebenso eine Energieform ist, gelingt es auch mit seiner Hilfe, die **Leitfähigkeit von Halbleitern beträchtlich zu erhöhen**. Deshalb werden Halbleiter-Bauelemente in Glasgehäusen mit einem lichtundurchlässigen — meistens schwarzen — Lack vor Lichteinwirkung geschützt. Die Steuerung elektronischer Schaltungen durch Licht weist erhebliche Vorzüge auf und wird in zunehmendem Maße ausgenutzt. Als wesentliche Vorzüge sind zu nennen: Das Licht legt in einer Sekunde 300 000 km zurück; die Wirkung des elektrischen Stroms hat nur bei **reinem Gleichstrom und sehr kurzen Leitungen** die gleiche Ausbreitungsgeschwindigkeit wie das Licht. Zur Übertragung von Licht werden keine Leitungen benötigt; es braucht keine Materie als Übertragungsmittel vorhanden zu sein. Licht ist praktisch trägeheitslos; man kann also beliebig schnelle Wechselvorgänge mit Hilfe des Lichts übertragen (z.B. Laserübertragung im GHz-Bereich).

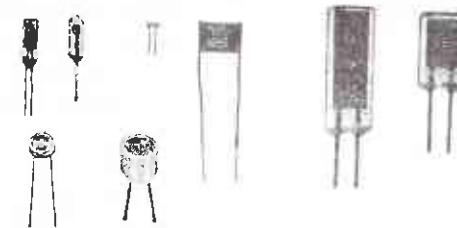


Abb. 2.25 -- Fotowiderstände

Als sehr brauchbare Halbleiterstoffe, deren Leitfähigkeit durch Lichteinwirkung stark erhöht wird, haben sich Cadmiumsulfid (CdS), Cadmiumselenid (CdSe) und Bleisulfid (PbS) erwiesen. Fotowiderstände enthalten meistens eine dünne mäanderförmige Schicht aus einem dieser lichtempfindlichen Halbleiterstoffe und sind vielfach in Glaskolben eingeschmolzen. Die lichtempfindliche Fläche hat eine Größe zwischen $0,01$ und 3 cm^2 ; einige Bauformen sind in Abb. 2.25 wiedergegeben (ca. $\frac{1}{3}$ natürl. Größe).

Als Schaltzeichen für Fotowiderstände wird das Symbol nach Abb. 2.26 verwendet. Die beiden entgegengerichteten Pfeilspitzen deuten an, daß es sich um einen stromrichtungsunabhängigen Widerstand handelt.

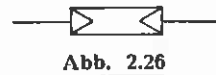


Abb. 2.26

2.2.4.1. Eigenschaften der Fotowiderstände

Im unbeleuchteten Zustand ist ein Fotowiderstand sehr hochohmig; sein sog. **Dunkelwiderstand** beträgt bis zu einigen Megohm. Schaltet man ihn in einen Gleichstromkreis, so fließt ein praktisch kaum meßbarer Strom. Mit zunehmender Beleuchtung fällt der Widerstand auf etwa 1/1000 des Dunkelwiderstands. Als Meßgröße des Lichts wird die Beleuchtungsstärke E in Lux zugrunde gelegt; z.B. ergibt eine Glühlampe 220 V/25 W in 1 m Abstand bzw. eine 100-W-Lampe in 2 m Abstand eine Beleuchtungsstärke von rd. 50 Lux. Die Abhängigkeit des Widerstands von der Beleuchtungsstärke ist im Diagramm Abb. 2.27 dargestellt.

Die größte Empfindlichkeit besitzen CdS-Fotowiderstände für rotes bis infrarotes Licht.

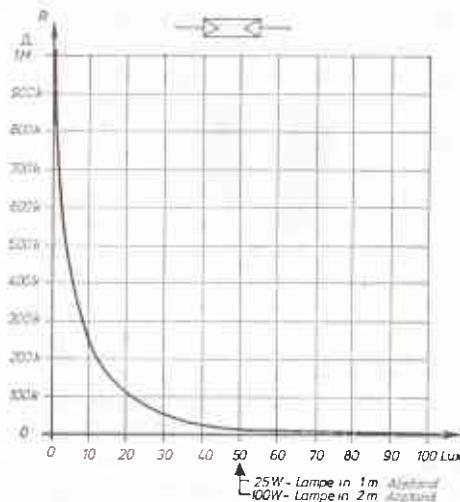


Abb. 2.27 — Widerstandsverhalten eines Fotowiderstands

Das Diagramm läßt erkennen, daß bereits verhältnismäßig geringe Beleuchtungsstärken ausreichen, um den Widerstand eines Fotowiderstands stark herabzusetzen oder, wie man auch sagt, den Fotowiderstand durchzusteuern. Im Gegensatz zu anderen fotoelektronischen Bauelementen sind Fotowiderstände verhältnismäßig hoch belastbar. Ihre zulässige Verlustleistung liegt etwa zwischen 0,05 und 1,2 Watt; sie eignen sich daher z.B. zur direkten Steuerung von Relaisstromkreisen. Eine einfache Versuchsschaltung für eine Lichtschranke mit einem Fotowiderstand ist in Abb. 2.28 wiedergegeben.

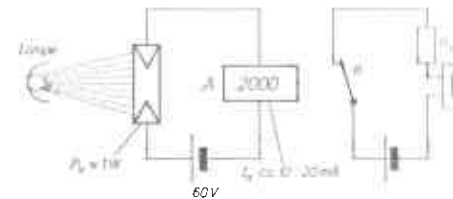


Abb. 2.28 — Einfache Lichtschranke mit einem Fotowiderstand

Der Fotowiderstand wird von einer Glühlampe beleuchtet. Zur Lichtbehandlung werden meistens Hohlspiegel (sog. Parabolspiegel) und Linsen eingesetzt. Das Relais A befindet sich so lange im angezogenen Zustand, wie Licht auf den Fotowiderstand fällt, der Ruhekontakt a ist somit geöffnet. Wird der Lichtstrahl

z.B. dadurch unterbrochen, daß jemand den Lichtstrahl durchschneidet, so erhöht sich der Widerstandswert des Fotowiderstands. Das Relais fällt ab, schließt den Ruhekontakt a und der Wecker ertönt.

Diese Art Lichtschranke kann zur Einbruchssicherung (natürlich mit unsichtbarem Infrarotlicht!), zum Einschalten von Rolltreppen, zur automatischen Öffnung von Türen oder auch zur elektrischen Zählung von Massenartikeln auf einem Fließband ausgenutzt werden. In Foto- und Filmapparaten dienen sogenannte CdS-Zellen zur automatischen Blendensteuerung. Dadurch entfällt die sonst notwendige Lichtmessung und Blendeneinstellung von Hand. Da Fotowiderstände zur Gruppe der passiven Bauelemente gehören, **muß in ihren Stromkreisen grundsätzlich eine Stromquelle (z.B. Batterie) vorhanden sein**. Leider arbeiten Fotowiderstände verhältnismäßig träge und reagieren nur auf Wechselspannung bis zu etwa 1000 Hz mit Widerstandsänderungen. Die im Abschnitt 3.6 behandelten Fotodioden sind dagegen für Frequenzen bis etwa 100 MHz geeignet.

3. Halbleiterdioden (Zweischicht-Halbleiter)

3.1. Grundsätzlicher Aufbau und Wirkungsweise einer Diode

Hauptmerkmal einer Halbleiterdiode ist das Vorhandensein eines PN-Übergangs; er ergibt sich an der Berührungs- oder Grenzschicht eines P-Leiters mit einem N-Leiter. Man kann einen PN-Übergang herstellen, indem ein Halbleiter-Kristall aus Germanium oder Silizium von einer Seite her N-dotiert und von der anderen Seite her P-dotiert oder indem ein N-Leiter von einer Seite her kräftig P-dotiert wird. Ein Bauelement, das nur einen PN-Übergang enthält, wird allgemein als Diode bezeichnet. An der Grenzschicht zwischen P- und N-Leiter stehen sich unmittelbar Löcher (bewegliche positive Ladungen) und Elektronen (bewegliche negative Ladungen) gegenüber.

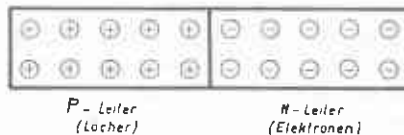


Abb. 3.1 — PN-Übergang

Ein PN-Übergang weist nun Eigenschaften auf, die ein einfaches aus einheitlichem Stoff aufgebautes Bauelement nicht besitzt. Wir wollen einmal untersuchen, was vorgeht, wenn man an einen PN-Übergang von außen eine Spannung anlegt.

3.1.1. Plus am P-Leiter; Minus am N-Leiter

Ursache für das Leitverhalten eines elektrisch leitfähigen Stoffs ist die elektrostatische Kraftwirkung: Gleichnamige elektrische Ladungen stoßen sich ab, ungleichnamige ziehen sich dagegen an. Mit Hilfe dieses einfachen physikalischen Gesetzes können wir das Verhalten nahezu aller Halbleiter-Bauelemente erklären. Was geht nun vor, wenn an einen PN-Übergang von außen eine Spannung mit Plus am P-Leiter und Minus am N-Leiter angelegt wird?

Die P-Ladungen (Löcher) werden vom Pluspol der Spannungsquelle abgestoßen und ebenso die N-Ladungen (Elektronen) vom Minuspol. Die beweglichen Ladungen werden somit über die Grenzschicht hinweg in das jeweils entgegengesetzt dotierte Gebiet abgedrängt.

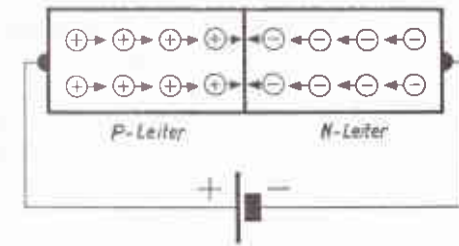


Abb. 3.2 — PN-Übergang in Durchlaßrichtung

So gelangen z.B. die Elektronen über das P-Gebiet zum Pluspol der Spannungsquelle; da Plus gleichbedeutend mit Elektronenmangel ist, findet ein Ausgleich statt, den wir üblicherweise als elektrischen Strom bezeichnen. Ebenso wandern die P-Ladungen (Löcher) in Richtung Minuspol; da Minus gleichbedeutend mit Elektronenüberschuß ist, nehmen die Löcher vom Minuspol Elektronen auf. Man kann also bei dieser Schaltung eines PN-Übergangs an sich nichts besonderes feststellen; denn er verhält sich ähnlich wie ein Leiter. In diesem Zustand ist der PN-Übergang stromdurchlässig.

Ein PN-Übergang ist stromdurchlässig geschaltet, wenn der Pluspol der äußeren Spannungsquelle am P-Leiter und der Minuspol am N-Leiter liegt. Dabei muß — wie wir noch sehen werden — die angelegte Spannung mindestens die Höhe der sogenannten Diffusionsspannung besitzen.

3.1.2. Minus am P-Leiter; Plus am N-Leiter

Die äußere Spannungsquelle wird umgepolzt. Was geht vor? Am P-Leiter liegt jetzt der Minuspol der Spannungsquelle. Da sich ungleichnamige Ladungen anziehen, werden die P-Ladungen (Löcher) zum Minuspol, also zum äußeren Anschluß hingezogen. Ebenso zieht der Pluspol die N-Ladungen (Elektronen) über den anderen äußeren Anschluß an. Die im Halbleiter vorhandenen beweglichen Ladungen wandern also von der Grenzschicht, dem eigentlichen PN-Übergang, ab. In der Nähe der Grenzschicht sind somit keine freien Ladungsträger mehr

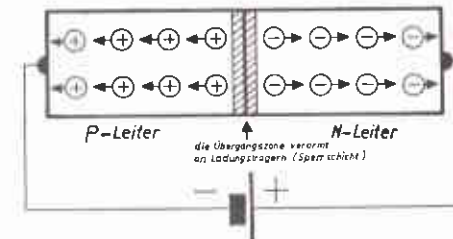


Abb. 3.3 — PN-Übergang in Sperrrichtung

vorhanden. Ein Stoff, der keine freien Ladungsträger besitzt, ist ein **Isolator**; über ihn kann kein Strom fließen. Unser PN-Übergang ist also stromundurchlässig geworden, er ist gesperrt.

Der gesperrte PN-Übergang ist mit einem Kondensator vergleichbar, dessen Dielektrikum die an Ladungsträgern verarmte Übergangszone darstellt. Diese Tatsache nutzt man in den sogenannten Kapazitätsdioden aus. Die an Ladungsträgern verarmte Übergangszone ist nur sehr dünn in der Größenordnung von einigen μm . Da jeder Isolierstoff bzw. jedes Dielektrikum eine bestimmte Durchschlagsfestigkeit aufweist, besteht für einen gesperrten PN-Übergang auch eine bestimmte Grenze für die Höhe der angelegten Spannung; diese Spannung wird als **Sperrspannung** bezeichnet.

Wird der Minuspol einer Spannungsquelle an den P-Leiter, der Pluspol an den N-Leiter eines PN-Übergangs gelegt, so fließt kein Strom; der PN-Übergang ist gesperrt. Mit Rücksicht auf die Durchschlagsfestigkeit darf die in Sperrrichtung angelegte Spannung bestimmte Höchstwerte nicht überschreiten. Eine Überschreitung der Sperrspannung führt fast immer zur Zerstörung einer Diode.

3.1.3. Die Diffusionsspannung

Ohne daß wir es beeinflussen könnten, spielt sich an der Grenzschicht eines PN-Übergangs ein Ladungsaustausch ab. Dieser Ladungsaustausch führt zur Entstehung der sogenannten Diffusionsspannung. Wie kommt es dazu?

P-Leiter und N-Leiter stehen sich an der Berührungsfläche unmittelbar gegenüber. Der P-Leiter hat das Bestreben, Elektronen einzufangen; der N-Leiter dagegen besitzt frei bewegliche Elektronen. Da die kleinen Bauteilchen der Materie bei Erwärmung in ständiger Bewegung sind, geraten die freien Elektronen infolge **Diffusion**¹⁾ an der Grenzschicht auch in die Nähe der Löcher und werden kurzerhand von diesen eingefangen; man nennt diesen Vorgang **Rekombinieren**²⁾. Mit diesem Hinüberwechseln von Elektronen aus dem N-Gebiet in das P-Gebiet findet aber gleichzeitig eine Ladungsverschiebung statt, denn Elektronen besitzen eine negative Ladung.

Vergleichen wir noch einmal: Ein P-Leiter ist ein Halbleiterstoff, der frei bewegliche positive Ladungen in Form der Löcher besitzt. Der N-Leiter dagegen enthält freie bewegliche negative Ladungen in Form der Elektronen. Beide sind jedoch von Natur aus nach außen hin **elektrisch neutral**, z.B. zeigt ein metallischer Leiter trotz der vielen freien Elektronen nach außen hin keine elektrische Ladung.

¹⁾ Diffusion = selbständige Vermischung an einer Berührungsschicht

²⁾ Rekombinieren = wieder zusammenfügen, den ursprünglichen Zustand wieder herstellen.

Wenn elektrische Ladungen von einem elektrisch neutralen Stoff entfernt und einem anderen elektrisch neutralen Stoff zugeführt werden, entsteht eine Spannung. Die an der Grenzschicht zwischen P- und N-Leiter durch Diffusion entstehende Spannung heißt **Diffusionsspannung** oder auch **Schwellspannung**. Sie ist aber nur in unmittelbarer Nähe der Grenzschicht wirksam und nicht etwa an den Außenelektroden mit einem Spannungsmesser nachweisbar.

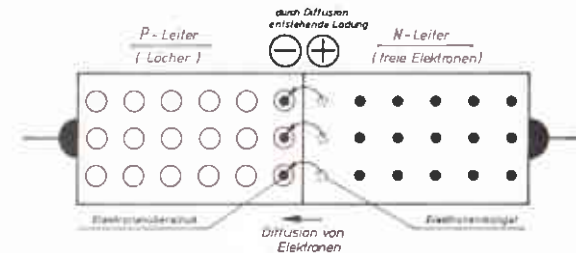


Abb. 3.4 — Entstehung der Diffusionsspannung

Sobald nämlich Elektronen über die Grenzschicht hinweg in die nahegelegenen Löcher gewandert sind, ist der N-Leiter wegen Elektronenmangels elektrisch positiv und versucht daher, die Elektronen zurückzuziehen. Jetzt stellt sich ein Gleichgewichtszustand ein zwischen der Kraft, mit der die Löcher die eingefangenen Elektronen festhalten und dem Bestreben des N-Leiters, die fehlenden Elektronen wieder zurückzuziehen. Die Kraftwirkung hängt von der Art des jeweiligen Halbleiterstoffs ab. Somit entsteht eine für den Halbleiterstoff typische Diffusionsspannung bestimmter Höhe. **Für PN-Übergänge ergeben sich bei den gebräuchlichen Halbleiterstoffen etwa folgende Diffusionsspannungen: Germanium: 0,2 V, Selen: 0,4 V, Silizium: 0,6 V.**

Da der P-Leiter an der Grenzschicht durch Elektronenaufnahme elektrisch negativ geladen ist, verhindert er ein weiteres Eindringen von Elektronen; denn nachfolgende Elektronen müßten ja gegen eine gleichnamige Ladung anlaufen. Das gleiche geht nun vor, wenn an einen PN-Übergang von außen eine Spannung in Durchlaßrichtung angelegt wird. Bei einem in Durchlaßrichtung geschalteten PN-Übergang müssen die freien Elektronen des N-Leiters in Richtung P-Leiter wandern. Daran werden sie jedoch gehindert, da der P-Leiter an der Grenzschicht eine negative Ladung aufweist. Erst wenn die äußere Spannung und mit ihr die elektrostatische Kraftwirkung größer wird als die der Diffusionsspannung, können die Elektronen des N-Leiters schließlich die Zone der negativen Ladung im P-Leiter überwinden. Die Höhe der Diffusionsspannung ist für einen bestimmten PN-Übergang nur dadurch meßbar, daß man ermittelt, bei welcher Höhe der Durchlaßspannung ein merklicher Stromfluß einsetzt.

Damit über einen PN-Übergang in Durchlaßrichtung ein Strom fließt, muß die angelegte Spannung größer als die Diffusionsspannung (Schwellspannung) sein.

3.1.4. Die Kennlinie einer Diode (PN-Übergang)

Am anschaulichsten wird das Verhalten eines Bauelements anhand seiner Spannungs-Strom-Kennlinie; sie zeigt, wie schon besprochen, die Abhängigkeit der Stromstärke von der Höhe der angelegten Spannung. Um die grundsätzliche Wirkungsweise einer Diode kennenzulernen, genügt die folgende Meßschaltung zur Kennlinienaufnahme einer Diode. Wegen der durch Spannungs- und Strommesser bedingten Fehler muß die Meßschaltung für die Durchlaß- bzw. Sperrrichtung geändert werden.

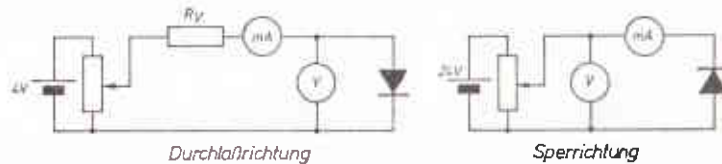


Abb. 3.5 — Meßschaltung zur Aufnahme einer Diodenkennlinie

Als Diode kann man eine Selenzelle aus einem Netzgleichrichter verwenden. Zum Schutz gegen Überlastung ist ein Schutzwiderstand R_V vorzusehen. Dieser ist so zu bemessen, daß der für die Diode höchstzulässige Strom nicht überschritten wird. Die Spannung wird zunächst in Durchlaßrichtung gepolt und mit Hilfe des Potentiometers stufenweise bis ca. 1,4 V erhöht. Zu jeder Spannung wird der

zugehörige Strom gemessen und anschließend die Meßreihe für den Sperrbereich aufgenommen (Höchstspannung etwa 22 bis 24 V). Die anhand der Meßwerte dargestellte Kennlinie zeigt etwa nebenstehenden Verlauf. Dazu ist noch zu bemerken, daß man aus rein praktischen Gründen für den Durchlaßbereich einer Diodenkennlinie einen anderen Maßstab für Span-

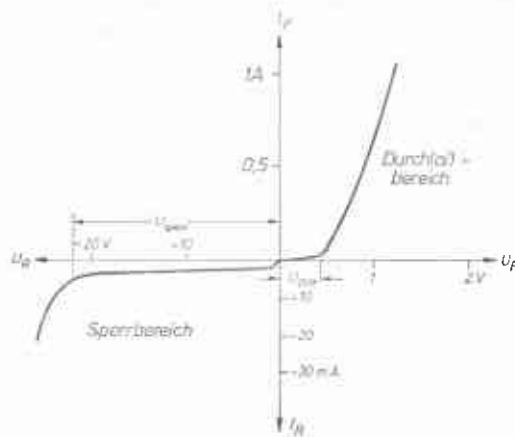


Abb. 3.6 — Kennlinie einer Diode (Selenzelle)

nung und Strom wählt als für den Sperrbereich. Die Spannung in Durchlaßrichtung wird mit U_F , der Durchlaßstrom mit I_F bezeichnet. Das F stammt aus dem Englischen: forward = vorwärts. Für Sperrspannung und Sperrstrom setzt man die Formelzeichen U_R und I_R , wobei R aus dem Englischen: return = zurückkehren, übernommen wurde. Bei der Diodenkennlinie ist grundsätzlich zwischen zwei Bereichen zu unterscheiden.

Durchlaßbereich: Solange die Spannung in Durchlaßrichtung kleiner als die Diffusionsspannung des PN-Übergangs ist, fließt ein kaum meßbarer Strom. Die Kennlinie verläuft noch sehr flach, was gleichbedeutend mit einem hohen Widerstand ist. Wird die Diffusionsspannung überschritten, so steigt der Strom stark an. Der Kennlinienverlauf wird steil, das läßt auf einen niedrigen Widerstandswert schließen.

Sperrbereich: Legt man an eine Diode eine Spannung in Sperrrichtung, so fließt ein — wenn auch geringer — Strom. Er wird durch die Eigenleitung des Halbleiterstoffs bei Raumtemperatur ermöglicht und steigt — wie bereits durch Versuch nachgewiesen — mit zunehmender Erwärmung. Bis zur Höhe der Sperrspannung verläuft die Kennlinie sehr flach. Die Diode ist also in diesem Bereich sehr hochohmig. Beim Erreichen der Sperrspannung beginnt jedoch der Durchbruch von Ladungsträgern; der Strom steigt auch im Sperrbereich an. Hohe Spannung und zunehmender Strom ergeben aber eine zunehmende Leistung = Verlustleistung; sie führt zu stärkerer Erwärmung und innerhalb kurzer Zeit zur Zerstörung der Diode. **Da die Spannungs-Strom-Kennlinie einer Diode nicht linear verläuft, ist ihr Widerstand nicht konstant. Er hängt von der Größe und Richtung der angelegten Spannung ab.** Diese Erkenntnis ist für viele Messungen an Dioden oder PN-Übergängen von großer Bedeutung für die richtige Beurteilung des Meßergebnisses.

Der Verlauf einer Diodenkennlinie läßt auch erkennen, daß man für das Verhalten der Diode keine mathematische Formel angeben kann. Eine genaue Beurteilung über das Verhalten einer Diode in einer bestimmten Schaltung ist daher nur mit Hilfe ihrer Kennlinie möglich.

3.1.5. Statischer und dynamischer Widerstand

Bevor wir uns mit praktischen Anwendungen und Messungen an Dioden beschäftigen, soll der Unterschied zwischen dem statischen und dynamischen Widerstand erläutert werden. In der Gleichstromlehre wird der Widerstand eines Bauelements mit Hilfe einer Span-

nungs- und Strommessung bestimmt. Der Widerstandswert läßt sich nach dem Ohmschen Gesetz berechnen ($R = U/I$). Eine Diode besitzt aber keinen bestimmten gleichbleibenden Widerstand. Zur besseren Beurteilung des Widerstandsverhaltens einer Diode hat man daher zwei Größen eingeführt, den statischen¹⁾ und den dynamischen²⁾ Widerstand. Wie die Bezeichnung schon erkennen läßt, bezieht sich der statische Widerstand auf einen „Ruhezustand“. Ein unveränderlicher — also statischer — elektrischer Zustand ergibt sich bei Gleichstrom. Legen wir an eine Diode in Durchlaßrichtung eine Gleichspannung U_F bestimmter Höhe, so fließt ein bestimmter Gleichstrom I_F . Aus diesen beiden Werten ergibt sich der

$$\text{statische Widerstand } R_F = \frac{U_F}{I_F}.$$

Er wird üblicherweise für den vorgesehenen Betriebszustand (Arbeitspunkt) einer Diode angegeben. Ebenso läßt sich für den Sperrbereich der statische Widerstand aus Sperrspannung U_R und Sperrstrom I_R bestimmen.

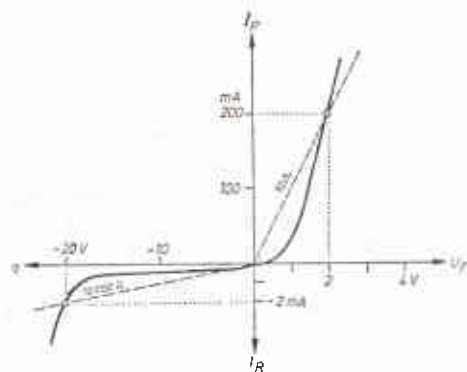


Abb. 3.7 — Bestimmung des statischen Widerstands einer Diode

Zeichnen wir für die beiden Widerstände $R_F = 10 \Omega$ und $R_R = 10\,000 \Omega$ die Widerstandsgeraden in das Kennlinienfeld, so ist deutlich zu erkennen, daß sie die Diodenkennlinie nur in je einem Punkt schneiden. Für alle anderen Spannungen und Ströme ergeben sich auch andere Diodenwiderstände.

¹⁾ statisch = in Ruhe

²⁾ dynamisch = in Bewegung

In der Diodenkennlinie nach Abb. 3.7 sind zwei Punkte markiert; für diese beiden Punkte ergeben sich folgende statische Widerstände:

$$\begin{aligned} R_F &= \frac{U_F}{I_F} \\ &= \frac{2 \text{ V}}{0,2 \text{ A}} = 10 \Omega, \\ R_R &= \frac{U_R}{I_R} \\ &= \frac{-20 \text{ V}}{-0,002 \text{ A}} = 10 \text{ k}\Omega. \end{aligned}$$

Anders liegen die Verhältnisse, wenn die Gleichspannung schwankt oder mit Wechselspannung überlagert wird. Angenommen, unsere Diode liegt an einer Gleichspannung von 2 Volt in Durchlaßrichtung.

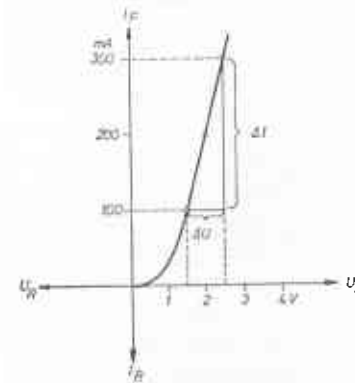


Abb. 3.8 — Bestimmung des dynamischen Widerstands einer Diode

Wird diese Gleichspannung mit einer Wechselspannung von $U_{\max} = 0,5$ Volt überlagert, so ändert sich die Spannung an der Diode im Rhythmus der Wechselspannung zwischen 1,5 und 2,5 Volt. Der Diodenstrom schwankt dabei zwischen 100 und 300 Milliampere. Für die Schwankungen — also den Wechselstrom — stellt die Diode einen Widerstand dar, der sich so berechnen läßt:

$$\begin{aligned} r &= \frac{\Delta U}{\Delta I} \\ &= \frac{2,5 \text{ V} - 1,5 \text{ V}}{0,3 \text{ A} - 0,1 \text{ A}} = \frac{1 \text{ V}}{0,2 \text{ A}} = 5 \Omega. \end{aligned}$$

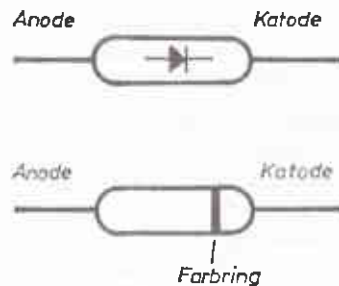
Dieser Widerstand ist für den angenommenen Arbeitspunkt also kleiner als der statische Widerstand; man bezeichnet ihn als **dynamischen** oder auch **differentiellen Widerstand**.

Der **statische Widerstand** eines Halbleiter-Bauelements ergibt sich aus **Gleichspannung** und **Gleichstrom**; er ist stark vom Arbeitspunkt abhängig. Unter dem **dynamischen Widerstand** r eines Halbleiter-Bauelements versteht man den Widerstandswert, der sich aus **Spannungsänderung** und **Stromänderung** ergibt; er stellt gewissermaßen einen Wechselstromwiderstand dar. Statischer und dynamischer Widerstand werden für einen Arbeitspunkt angegeben. Der dynamische Widerstand wird ganz einfach dadurch ermittelt, daß man im Arbeitspunkt eine Tangente an die Kennlinie legt. Die Tangente stellt dann die Widerstandsgerade für den dynamischen Widerstand dar.

Unter dem **Arbeitspunkt** ist allgemein der Punkt auf einer Kennlinie zu verstehen, der durch eine bestimmte Gleichspannung oder einen bestimmten Gleichstrom festgelegt wird. Er wird so gewählt, daß das Bauelement für den vorgesehenen Zweck einwandfrei arbeitet, ohne überlastet zu werden. Bei der Anwendung von Wechselspannungen muß der Arbeitspunkt so liegen, daß der positive und der negative Scheitelwert der Wechselspannung im geradlinigen Teil der Kennlinie liegen, um Verzerrungen zu vermeiden.

3.2. Messungen an Dioden

Eine Diode bietet in der Elektrotechnik und Elektronik eine Vielzahl von Anwendungsmöglichkeiten. Wird sie nicht speziell für einen ganz besonderen Anwendungszweck hergestellt, so heißt sie Universaldiode; als Schaltzeichen wird nebenstehendes Symbol verwendet. Dabei ist wichtig zu wissen, daß die Pfeilspitze des Symbols immer in die Durchlaßrichtung (technische Stromrichtung!) des PN-Übergangs weist. Dioden gehören zu den passiven Bauelementen, also zu den „Verbrauchern“ elektrischer Energie. Über einen Verbraucher fließt der Strom von Plus nach Minus bzw. von der Anode zur Katode. Sinngemäß werden die beiden Anschlüsse einer Diode für die Durchlaßrichtung mit Anode und Katode bezeichnet. Zur Kennzeichnung wird bei kleinen Bauformen allgemein die Anschlußseite der Katode mit einem weißen oder roten Farbring oder auch Farbpunkt versehen; auf größere Diodengehäuse wird einfach das Diodensymbol gedruckt. Die Katode einer Diode ist immer die Stromaustrittsstelle für die Durchlaßrichtung; Kennzeichnung und Bauformen sind aus Abb. 3.10 ersichtlich (weitere Bauformen sind im Anhang — Abschnitt 10 — wiedergegeben).



Universaldiode (Glasgehäuse)
($P_V \approx 100 \text{ mW}$)



Leistungsdiode (bis ca. 1 A)
($P_V \approx 1 \text{ W}$)



Abb. 3.10 — Anschluß-Kennzeichnung und Bauformen von Dioden



Abb. 3.9

3.2.1. Einfache Messungen mit dem Ohmmeter

Normalerweise sind die Anschlüsse handelsüblicher Dioden gekennzeichnet. In der Praxis kommt es jedoch vor, daß Farbringe oder -punkte verwischt sind oder die Diodenanschlüsse Anode und Katode äußerlich nicht erkennbar sind. Hier hilft uns in einfacher Weise ein Ohmmeter.

Die Anschlußbuchsen bzw. Meßschnüre eines Ohmmeters sind mit $\oplus$ und $\ominus$ gekennzeichnet. Die im Ohmmeter vorhandene Batterie (meistens eine Einzelzelle 1,5 V) liegt mit ihrem Pluspol an der $\oplus$ -Buchse des Ohmmeters. Ihre Spannung ist größer als die Diffusions- oder Schwellspannung der Dioden (Ge = 0,2 V; Se = 0,4 V; Si = 0,6 V).

Legen wir eine Diode an die Meßschnüre eines Ohmmeters und vertauschen anschließend die beiden Meßschnüre, so wird die Diode mit der im Ohmmeter eingebauten Spannungsquelle einmal in Durchlaßrichtung und zum anderen in Sperrichtung betrieben. Den Durchlaßbereich erkennen wir an einem niedrig angezeigten Widerstandswert, die Sperrichtung an einem hohen Widerstandswert. Bei Anzeig eines niedrigen Widerstandswerts ist der mit der $\ominus$ -Buchse bzw. $\ominus$ -Meßschnur des Ohmmeters verbundene Diodenanschluß die Katode.

Werden die Messungen mit einem Vielfachinstrument durchgeführt, so ist zu beachten, daß die Polarität dann umgekehrt ist! Bei Messungen an Halbleiter-Bauelementen mit einem Vielfachinstrument empfiehlt sich, die rot gekennzeichnete Meßschnur an die $\ominus$ -Buchse des Vielfachinstruments anzuschließen. Da in Gleichstromnetzen der Plusleiter immer rot gekennzeichnet wird, brauchen wir uns jetzt bei Messungen nur zu merken:

Rot $\triangleq \oplus$.

kleiner Widerstandswert

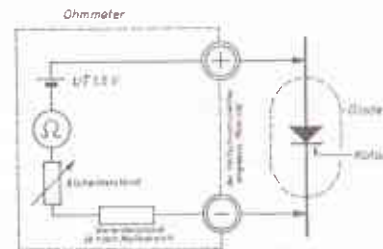
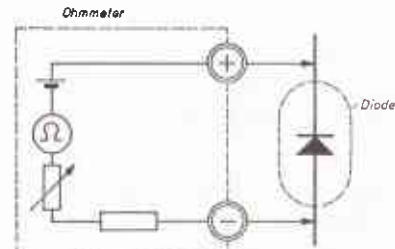


Abb. 3.11 — Durchlaßrichtung

Zur Messung wird der Widerstandsbe-
reich $\times 1k$ eingeschaltet. Bei diesem Meß-
bereich liegt in der Skalenmitte ein Wert
von ca. 30...50 k Ω . Zur Messung benut-
zen wir die Universaldiode AA 118 und
anschließend eine Siliziumdiode, z.B. BA
147. (Über die Bezeichnung der Halbleiter-
Bauelemente lesen Sie bitte im Anhang
dieses Buches nach.)

Für die Diode AA 118 ergeben sich etwa
folgende Meßwerte:

- a) Durchlaßbereich = 2,5 k Ω ,
- b) Sperrbereich = 2 M Ω .

großer Widerstandswert**Abb. 3.12 — Sperrrichtung**

Schalten wir für den Durchlaßbereich den Ohmmeter-Meßbereich auf $\times 10$ um, so ergibt sich ein Durchlaßwiderstand von rd. 200 Ohm.

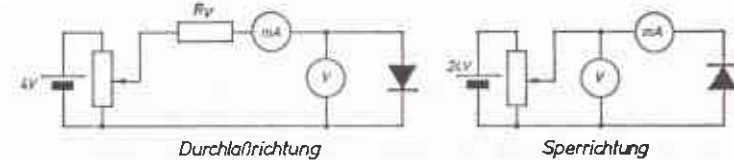
Daß sich für den Durchlaßbereich sehr unterschiedliche Widerstandswerte bei Umschaltung des Meßbereichs ergeben, hängt damit zusammen, daß der Innenwiderstand des Ohmmeters bei einer Meßbereichsumschaltung ebenfalls geändert wird. Infolge der geänderten Reihenschaltung Innenwiderstand — Diode ändert sich aber auch die Teilspannung an der Diode. Von der Kennlinie einer Diode her wissen wir, daß der statische Widerstand einer Diode sehr stark von der anliegenden Spannung abhängig ist.

Aus den gemessenen Widerstandswerten für den Durchlaß- und Sperrbereich kann man aber auch — zumindest grob — auf die Funktionsfähigkeit der Diode schließen. Dafür gilt folgende Richtlinie: **Bei einer einwandfreien Diode sollte das Verhältnis von gemessenem Sperrwiderstand zu gemessenem Durchlaßwiderstand (im gleichen Widerstandsmeßbereich) etwa 100 oder mehr betragen.** Dieser Wert ist rein größenordnungsmäßig zu verstehen. Für Si-Dioden muß dieses Verhältnis größer sein als für Ge-Dioden, da Si-Dioden einen kleineren Durchlaß- und einen größeren Sperrwiderstand aufweisen.

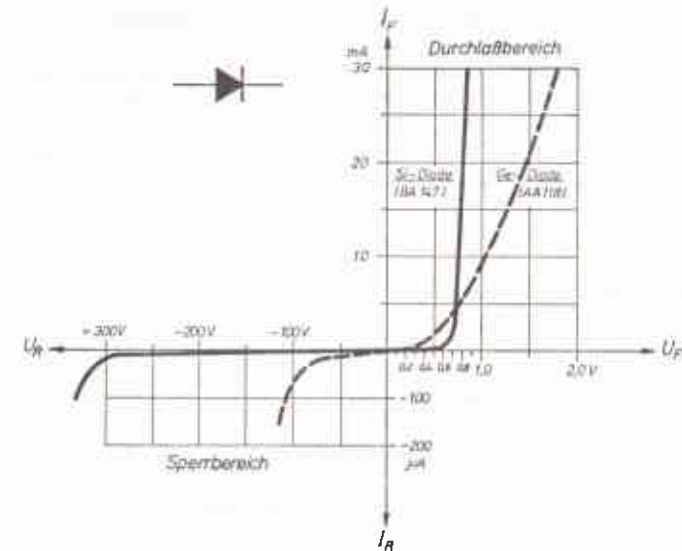
Die Messungen mit einem Ohmmeter lassen erkennen, daß man Dioden auch mit einfachen Mitteln verhältnismäßig schnell prüfen kann. Wir werden solche Messungen später auch an Transistoren vornehmen.

3.2.2. Die Kennlinien verschiedener Dioden

Um einen grundsätzlichen Vergleich zwischen Dioden aus den verschiedenen Halbleiterstoffen zu ermöglichen, sollen die Kennlinien einer Germaniumdiode (AA 118) und einer Siliziumdiode (BA 147) aufgenommen und in einem Diagramm dargestellt werden. Da die Sperrspannung für Germanium- und vor allem Siliziumdioden schon verhältnismäßig hoch liegt, muß auf die versuchsmäßige Aufnahme der Sperrkennlinien bis zum Durchbruch verzichtet werden. Die Sperrkennlinien für die gemessenen Dioden wurden nach Herstellerangaben in das Kennlinienfeld eingetragen. Für den Sperrbereich genügt eine Messung bis etwa 50 V. Zur Kennlinienaufnahme können wir uns der schon bekannten Meßschaltungen bedienen.

**Abb. 3.13 — Meßschaltung zur Kennlinienaufnahme von Dioden**

Wegen des Überganges von hochohmigem zu niederohmigem Verhalten im Durchlaßbereich ergeben sich genaue genommen Meßfehler für den Bereich unterhalb der Diffusionsspannung. Da jedoch beide Dioden in der gleichen Meßschaltung untersucht werden, ist trotz dieser Meßfehler ein Vergleich möglich. Zudem machen sich Meßfehler unterhalb der Diffusionsspannung wegen des größeren Maßstabs im Durchlaßbereich der Kennlinie praktisch nicht bemerkbar. Zeichnet man die Kennlinien, so ist folgendes zu erkennen:

**Abb. 3.14 — Kennlinien einer Germanium- und einer Siliziumdiode**

- Die Kennlinie der Siliziumdiode verläuft im Durchlaßbereich oberhalb der Diffusionsspannung steiler als bei der Germaniumdiode. Die Siliziumdiode besitzt also im Vergleich zur Germaniumdiode den kleineren dynamischen Durchlaßwiderstand.
- Der Kennlinienknick im Durchlaßbereich läßt auf die Diffusionsspannungen schließen (ca. 0,2 V für Ge und ca. 0,6 V für Si). Für bestimmte Aufgaben ist die hohe Diffusionsspannung von

Si-Dioden ungünstig. Daher sind Si-Dioden z.B. nicht als Meßgleichrichter oder Hochfrequenz-Gleichrichter geeignet.

- c) Im Sperrzustand fließt über die Dioden nur ein sehr geringer Sperrstrom in der Größenordnung von μA , bei der Siliziumdiode sogar nur in der Größenordnung von nA . Das beweist, daß der Sperrwiderstand von Si-Dioden um mehrere Größenordnungen größer sein kann als der von Ge-Dioden. An den Sperrkennlinien zeigt sich auch deutlich der Übergang zum Ladungsträgerdurchbruch in der Nähe der Sperrspannung.
- d) Kommt es auf eine hohe Sperrspannung an, so ist die Siliziumdiode allen anderen Halbleiter-Gleichrichtern weit überlegen.

Da Selenzellen meistens nur aus Leistungsgleichrichtern verfügbar sind, ist ein unmittelbarer Vergleich in dem für Kleinleistungsdioden dargestellten Kennlinienfeld nicht möglich. Die wichtigsten Eigenschaften der drei Gleichrichterarten sind in folgender Tabelle zusammengestellt:

	Sperrspannung V	Diffusionsspannung V	Strombelastbarkeit $\frac{\text{A}}{\text{cm}^2}$	Zul. Grenztemperatur $^{\circ}\text{C}$	Wirkungsgrad %	Verwendung als
Selenzelle	20...30	0,45	0,1	80	80	Netzgleichrichter, Gegenzelle, Gehörschutzgleichrichter
Germaniumdiode	30...120	0,20	80	75	90	HF-Gleichrichter, Begrenzer, Schalter, Meßgleichrichter
Siliziumdiode	100...2000	0,65	200	150	99,6	Leistungsgleichrichter, Begrenzer, Schalter, Kapazitätsdiode

3.3. Universaldioden und ihre Anwendung

3.3.1. Die Diode als Gleichrichter

3.3.1.1. Einweggleichrichter

Wegen ihres ausgeprägten Durchlaß- oder Sperrverhaltens wird die Diode zur Gleichrichtung von Wechselströmen eingesetzt. Schalten wir in einen Wechselstromkreis eine Diode in Reihe zum Verbraucher,

so kann der Strom immer nur in einer Richtung — nämlich in Durchlaßrichtung der Diode — fließen; denn für die entgegengesetzte Stromrichtung ist die Diode so hochohmig, daß praktisch ein kaum meßbarer Strom fließt.

Die einfachste Form des Wechselstroms ist der sinusförmige Wechselstrom; d.h., das Strom-Zeit-Diagramm stellt eine Sinuskurve dar. Im

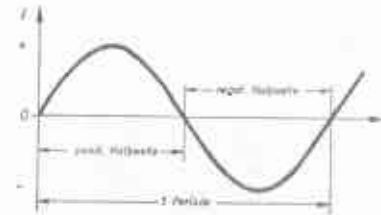


Abb. 3.15 — Zeitlicher Verlauf des Sinuswechselstroms

Stromkreis mit eingeschalteter Diode kann aber von den beiden Halbwellen des Wechselstroms nur eine wirklich über den Verbraucher fließen. Während der anderen Halbwelle ist die Diode gesperrt. Wenn aber kein Strom über den Verbraucherwiderstand fließt, dann tritt an ihm auch kein Spannungsabfall auf. Wir erkennen daraus, daß der Verbraucherwiderstand R nur während einer Halbwelle Strom bzw. Spannung erhält. Wegen dieser Eigenschaft bezeichnet man die einfachste Gleichrichterschaltung als Einweggleichrichter-Schaltung oder kurz Einwegschaltung.

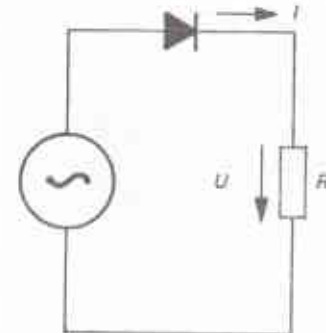


Abb. 3.16 — Einwegschaltung

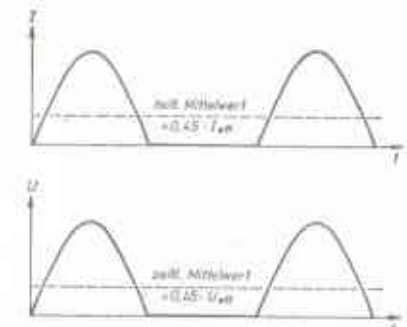


Abb. 3.17 — Strom und Spannung an R

In Abb. 3.17 ist der zeitliche Verlauf des Gleichstroms dargestellt, der durch eine Einweggleichrichterschaltung aus sinusförmigem Wechselstrom gewonnen wurde. Da der Strom an R einen Spannungsabfall ohne Phasenverschiebung verursacht, hat die Spannung an R den gleichen zeitlichen Verlauf. Man nennt diese Stromart in der Stromversorgungstechnik pulsierenden Gleichstrom. Zur Darstellung des zeitlichen Verlaufs bietet sich das Oszilloskop an; dazu wollen wir die folgende Schaltung aufbauen.

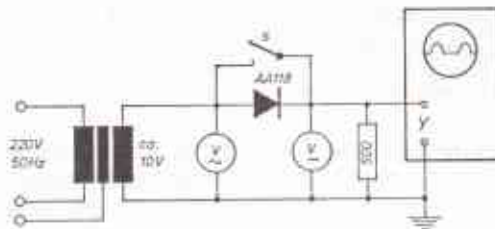


Abb. 3.18 — Darstellung des Stromoszillogramms für eine Einweggleichrichter-Schaltung

Der Y-Eingang des Oszilloskops ist auf = zu schalten. Ist der Schalter *s* geschlossen, so zeigt sich auf dem Bildschirm ein sinusförmiges Wechselstromdiagramm. Beim Öffnen des Schalters fallen die negativen Halbwellen fort; deutlich ist die Übereinstimmung mit dem in Abb. 3.17 dargestellten Diagramm zu erkennen. Von Gleichstrom kann aber noch keine Rede sein. Man spricht daher von pulsierendem Gleichstrom.

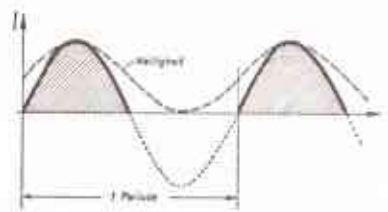


Abb. 3.19 — Pulsierender Gleichstrom nach Einweggleichrichtung

In erster Annäherung kann man die Frequenz der Welligkeit mit der Frequenz des gleichgerichteten Wechselstroms gleichsetzen; d.h., ein durch Einwegschaltung gleichgerichteter 50-Hz-Wechselstrom erzeugt in einem Kopfhörer oder Lautsprecher einen tiefen 50-Hz-Ton.

Der zeitliche Mittelwert der pulsierenden Gleichspannung ist bei Einweggleichrichtung verhältnismäßig gering; er beträgt nur 45 % vom Effektivwert der Wechselspannung. Das bedeutet, daß — abgesehen von den Spannungsverlusten an der Diode — ein Spannungsmesser vor der Diode 10 V Wechselspannung, hinter der Diode dagegen nur 4,5 V Gleichspannung anzeigen würde; das ist in der Meßschaltung mit Hilfe der beiden Spannungsmesser nachweisbar.

Ein wesentlicher Nachteil der Einwegschaltung ist, daß nur je eine Halbperiode des Wechselstroms ausgenutzt wird. Um aus pulsierendem Gleichstrom reinen Gleichstrom (gleichbleibende Größe) zu gewinnen, muß man hinter den Gleichrichter eine Siebkette schalten.

3.3.1.2. Zweiweggleichrichter

In der Zweiweggleichrichter-Schaltung (auch kurz Zweiwegschaltung genannt) werden beide Halbwellen des Wechselstroms ausgenutzt. Die Schaltung wurde früher überwiegend für Netzgleichrichter mit Röhren angewandt. Heute gewinnt sie in der Fernmeldetechnik z.B. zur Frequenzverdoppelung wieder an Bedeutung.

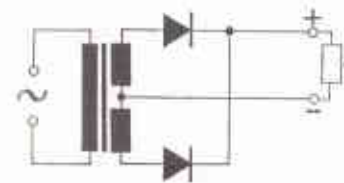


Abb. 3.20 — Zweiweggleichrichter-Schaltung

Für die Zweiwegschaltung wird ein Transformator mit zwei in Reihe geschalteten Sekundärwicklungen benötigt. Wegen dieses Aufwands hat die Zweiwegschaltung in Stromversorgungsanlagen keine Bedeutung mehr. Betrachtet man in den beiden Abbildungen 3.21 links und rechts den Zustand der Gleichrichterdioden, so ergibt sich:

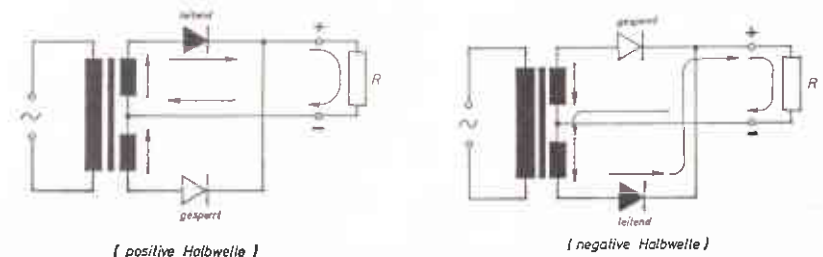


Abb. 3.21 — Zweiwegschaltung bei positiver und negativer Halbwellen

werden aber jetzt beide Halbwellen des Wechselstroms in gleicher Richtung über den Verbraucher geleitet. Der Verbraucher erhält pulsierenden Gleichstrom, dessen zeitlicher Verlauf in Abbildung 3.22 wiedergegeben ist. In erster Annäherung ergibt sich für die Welligkeit des pulsierenden Gleichstroms die doppelte Frequenz gegenüber dem zugeführten Wechselstrom. Ein nachgeschalteter Kopfhörer würde bei Gleichrichtung eines 50-Hz-Wechselstroms einen 100-Hz-Ton wiedergeben (1 Oktave höher).

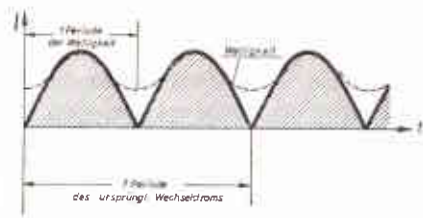


Abb. 3.22 — Zeitlicher Verlauf des Stroms nach Zweiweggleichrichtung

Diese Eigenschaft einer Zweiwegschaltung macht man sich in der Fernmeldetechnik zur Frequenzverdoppelung zunutze. So wird z.B. in Stereo-Rundfunkempfängern aus dem vom Sender ausgestrahlten 19-kHz-Stereo-Pilotton mit Hilfe einer Zweiweggleichrichter-Schaltung die zum Stereoempfang notwendige 38-kHz-Wechselspannung gewonnen.

3.3.1.3. Brückengleichrichter

Die in Abbildung 3.23 wiedergegebene Brückenschaltung verbindet die Vorteile eines einfachen Aufbaus mit der Ausnutzung beider Halbwellen des Wechselstroms; denn hier ist kein Transformator erforderlich. Die Brückenschaltung wird daher zur Wechselstrom-Gleichrichtung am häufigsten angewandt. Sie besteht aus vier zu einer Brücke zusammengeschalteten Dioden.

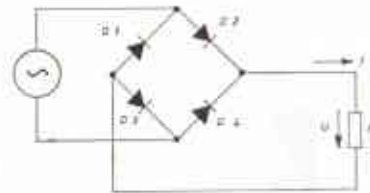


Abb. 3.23 — Brückenschaltung

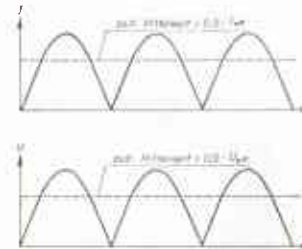


Abb. 3.24 — Strom und Spannung an R

Die Wirkungsweise läßt sich am einfachsten erklären, wenn man die Schaltung jeweils bei positiver und negativer Halbwelle des zugeführten Wechselstroms betrachtet. Wir wollen annehmen, daß bei der positiven Halbwelle der Strom vom oberen Anschluß der Wechselstromquelle abfließt. Für diese Stromrichtung sind aber nur die Dioden D2 und D3 durchlässig, während die Dioden D1 und D4 sperren. Die gesperrten Dioden sind jeweils gestrichelt dargestellt; der Stromkreis ist über D2 — R — D3 geschlossen.

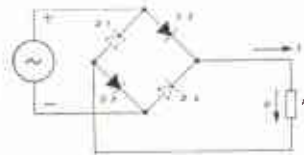


Abb. 3.25 — Strom bei positiver Halbwelle

Kehrt die Stromrichtung bei der negativen Halbwelle um, so sind die Dioden D1 und D4 in Durchlaßrichtung, die Dioden D2 und D3 in Sperrrichtung geschaltet. Der Stromkreis ist jetzt über D4 — R — D1 geschlossen. Beim Vergleich der Stromrichtungen über R bei positiver und negativer Halbwelle ist festzustellen, daß sie in beiden Fällen gleich sind. Über den Verbraucherwiderstand fließt der Strom also immer in gleicher Richtung; er ist jedoch nicht zeitlich konstant.

Abb. 3.26 — Strom bei negativer Halbwelle

Den Nachweis für die Richtigkeit dieser Überlegungen führen wir wieder mit dem Oszilloskop durch.

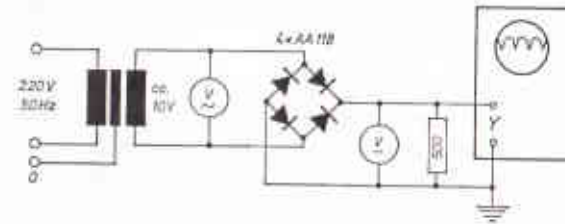


Abb. 3.27 — Darstellung des Stromoszillogramms für eine Brückengleichrichter-Schaltung

Auch hier zeigt sich volle Übereinstimmung mit dem in Abb. 3.24 dargestellten Diagramm. Da bei der Brückengleichrichtung beide Halbwellen des zugeführten Wechselstroms ausgenutzt werden, ist der zeitliche Mittelwert der am Verbraucher R liegenden pulsierenden Gleichspannung entsprechend doppelt so groß wie bei der Einweggleichrichtung. Er beträgt 90 % vom Effektivwert der angelegten Wechselspannung, was an den eingeschalteten Spannungsmessern abzulesen ist.

3.3.1.4. Siebung

Aus allen Gleichrichterschaltungen wird, wie wir gesehen haben, pulsierender Gleichstrom gewonnen. Zum Betrieb vieler Geräte, z.B. Verstärker, ist jedoch nur reiner Gleichstrom geeignet, da pulsierender Gleichstrom einen starken Brummtton verursachen würde. Um reinen Gleichstrom zu erhalten, muß der pulsierende Gleichstrom einer Gleichstromrichterschaltung „geglättet“ werden. Dazu müssen die Welligkeitsanteile beseitigt, d.h. ausgesiebt werden. Das wird durch Siebketten erreicht.

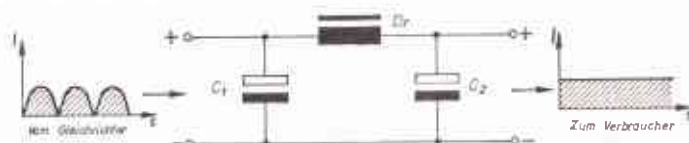


Abb. 3.28 — Siebkette

Siebketten bestehen aus Kondensatoren und Drosselspulen oder in vereinfachten Schaltungen auch aus Kondensatoren und Widerständen. Der erste Kondensator einer Siebkette liegt parallel zum Ausgang der Gleichrichterschaltung. Er wird Ladekondensator genannt und beeinflusst die Spannung. Steigt die Spannung, so wird der Kondensator aufgeladen; er nimmt Energie auf. Diese gespeicherte Energie gibt er wieder an den Verbraucher ab, wenn die Spannung aus der Gleichrichterschaltung absinkt. **Der Kondensator wirkt also während der Spannungslücken als Spannungsquelle.**

Die Drosselspule liegt in Reihe mit dem Verbraucher. Sie beeinflusst den Strom. Nach dem Lenzschen Gesetz über den Induktionsvorgang wird bei einer Induktivität (Drossel) beim Einschalten — also bei Stromanstieg — ein Induktionsstrom entgegengesetzter Richtung erzeugt. Beim Ausschalten — also bei Stromabfall — hat der in der Induktivität erzeugte Selbstinduktionsstrom die gleiche Richtung wie der zugeführte Strom. **Die Wirkung der Drosselspule ist also bei Stromanstieg stromhemmend, bei Stromabfall stromerhaltend. Somit werden die Stromspitzen abgeflacht und die „Stromtäler“ ausgefüllt.**

Der zweite Kondensator glättet wiederum die Spannung. Bei ausreichender Bemessung der Kapazitäten und Induktivität ist hinter einer Siebkette praktisch reiner Gleichstrom vorhanden. **Je größer Induktivität und Kapazitäten sind, desto besser ist die Siebwirkung einer Siebkette. Bei nur geringem Strombedarf kann die Drosselspule durch einen ohmschen Widerstand ersetzt werden.**

3.3.2. Die Diode als Schalter

Wir haben anhand der Kennlinie einer Halbleiterdiode zwei deutlich verschiedene Zustände in ihrem Verhalten feststellen können:

- Die in Durchlaßrichtung geschaltete Diode wird niederohmig, sobald die anliegende Spannung den Wert der Diffusionsspannung geringfügig überschreitet. Eine Siliziumdiode geht also oberhalb von etwa 0,6 Volt, eine Germaniumdiode schon oberhalb von etwa 0,2 Volt in den niederohmigen Zustand über.

- Ist die Spannung, die an einer in Durchlaßrichtung geschalteten Diode liegt, kleiner als die Diffusionsspannung oder wird sie umgepolt, so ist die Diode hochohmig.

Man kann die beiden widerstandsmäßig sehr unterschiedlichen Zustände einer Diode mit denen eines mechanischen Schalters vergleichen: Ein geschlossener Schalter ist — wie die in Durchlaßrichtung geschaltete Diode — niederohmig. Andererseits ist der geöffnete Schalter ähnlich wie eine gesperrte Diode äußerst hochohmig. Eine Diode kann somit als Schalter eingesetzt werden.

Der Vergleich zwischen mechanischem Schalter und Diode ist jedoch etwas unkorrekt. Bei genauer Untersuchung müssen wir nämlich feststellen, daß der Widerstand eines geschlossenen Schalters in der Größenordnung von tausendstel Ohm, der Widerstand einer durchlässig geschalteten Diode immerhin noch in der Größenordnung Ohm liegt. Der geöffnete mechanische Schalter hat einen Widerstand von praktisch unendlich Ohm, während die Diode einen Sperrwiderstand in der Größenordnung Megohm besitzt. Eine Diode ist also bezüglich ihres Durchlaß- und Sperrwiderstands kein idealer Schalter. Sie weist aber gegenüber mechanischen und elektromechanischen Schaltern Vorzüge auf, die zu einer inzwischen breiten Anwendung in der Elektrotechnik und Elektronik geführt haben. Zu den Vorteilen einer Diode als Schalter (kurz als Schalterdiode bezeichnet) gegenüber einem mechanischen Schalter gehört die wesentlich kürzere Ein- und Ausschalzeit. Mit Dioden erreicht man Schaltzeiten ¹⁾ von weniger als 100 ns (1 ns = 1 Nanosekunde = 1/1 000 000 000 Sekunde), während die Schaltzeit mechanischer Schalter in der Größenordnung von ms liegt (1 ms = 1 Millisekunde = 1/1 000 Sekunde).

Bei mechanischen Schaltern können Prellungen auftreten, die auf der Massenträgheit und Elastizität der Kontakte, Kontaktarme und Kontaktfedern beruhen. Prellungen wirken sich dahingehend aus, daß der Kontakt beim Schließen noch ein oder mehrere Male zurückprellt und nicht gleich fest geschlossen ist.

Insbesondere die kurzen Schaltzeiten und die Prellfreiheit der Schalterdioden ermöglichen in der Datenverarbeitung und elektronischen Vermittlungstechnik eine ganz erhebliche Verkürzung des Funktionsablaufs und Verbesserung der Betriebssicherheit.

¹⁾ Schaltzeit ist die Zeit, die zwischen Auslösung und endgültiger Schalterstellung liegt.

Wir wollen die Wirkungsweise wieder an einer einfachen Schaltung mit Dioden und Glühlampen untersuchen.



Abb. 3.29 — Dioden als Schalter

Steht der Polwechschler in Schalterstellung 1, so liegt an der oberen Leitung das Pluspotential der Spannungsquelle, an der unteren das Minuspotential. Stromdurchlässig ist jetzt nur die Diode D_1 ; deshalb leuchtet auch nur die Signallampe 1 auf; die Diode D_2 ist gesperrt. Wird der Polwechschler in Stellung 2 umgeschaltet, so ist die Diode D_2 stromdurchlässig, während die Diode 1 in Sperrrichtung betrieben wird. Darum leuchtet jetzt die Signallampe 2.

In der Schaltungspraxis wird die Steuerung nicht — wie in der vereinfachten Versuchsschaltung — durch mechanische Schalter, sondern durch Rechteckwechselspannungen vorgenommen.

In der folgenden Schaltung dient eine Diode als Wechselstromschalter. Die Diode ist parallel zum Verbraucher geschaltet und schließt ihn im Schaltzustand EIN kurz. Der als Verbraucher dienende Fernhörer ist über die Kondensatoren C_1 und C_2 sowie einen Schutzwiderstand R_V mit einem Tongenerator verbunden. Parallel zum Hörer ist eine Diode geschaltet. Über zwei besondere Zuleitungen kann mit Hilfe eines Polwechschlers eine Spannung mit wechselnder Polarität angelegt werden. In einer dieser Leitungen liegt ebenfalls ein Schutzwiderstand R_S . Die Schutzwiderstände sollen verhindern, daß die Gleichspannungsquelle oder der Tongenerator kurzgeschlossen bzw. die Diode überlastet wird.

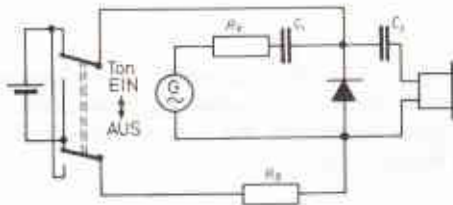


Abb. 3.30 — Diode als Wechselstromschalter

Steht der Polwechschler auf EIN, so wird die Diode gesperrt. Somit ist im Fernhörer deutlich ein Ton hörbar. Beim Umschalten liegt die Diode für die Gleichspannung in Durchlaßrichtung. Sie ist niederohmig und schließt den Fernhörer kurz. Die Lautstärke, mit der der Ton wahrnehmbar ist, geht deutlich zurück. Er verschwindet jedoch nicht vollständig, weil der Durchlaßwiderstand der Diode nicht Null wird.

Wir können mit Hilfe dieser Schaltung den Fernhörer über beliebig lange Leitungen sehr einfach ein- und ausschalten. Im Prinzip entspricht diese Schaltung einem Modulator ¹⁾, wie er in der Trägerfrequenztechnik verwendet wird.

Dioden können als Schalter in Reihe zum „Verbraucher“ geschaltet werden und haben die Funktion eines üblichen EIN-AUS-Schalters. Sie können aber auch parallel zum Verbraucher geschaltet werden und schließen diesen im Durchlaßzustand kurz. Infolge der Nebenschlußschaltung ergibt sich, daß der Verbraucher außer Betrieb ist, wenn der Diodenschalter geschlossen ist und daß er sich im Betriebszustand befindet, wenn die Diode gesperrt — also als Schalter geöffnet ist. Damit ergibt sich ein im Vergleich zu einem üblichen Schalter entgegengesetztes Verhalten. Man spricht auch von einer Umkehrschaltung.

3.3.3. Die Diode als Entkoppler

Die Ventilwirkung von Dioden kann auch zur Entkopplung von Stromkreisen ausgenutzt werden. Dabei soll verhindert werden, daß trotz galvanischer Verbindung unerwünscht Strom in benachbarte Stromkreise abfließen kann. Zur Veranschaulichung soll eine Schaltung mit zwei Relais, zwei Dioden und drei Schaltern dienen. Die Schaltung soll folgende Funktionen ermöglichen: Die beiden Relais A und B sollen mit Hilfe der drei Schalter 1, 2 und 3 so gesteuert werden, daß der Schalter 1 nur das Relais A, der Schalter 2 beide Relais und der Schalter 3 nur das Relais B einschaltet.

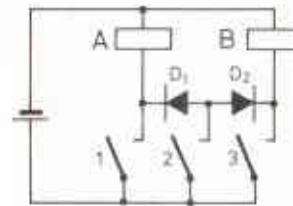


Abb. 3.31 — Dioden
als Entkoppler

In der dargestellten Schaltung verhindert die Diode D_1 bei geschlossenem Schalter 1 den Anzug des B-Relais und sinngemäß die Diode D_2 bei geschlossenem Schalter 3 den Anzug des A-Relais. Beide Dioden sind nur dann gleichzeitig leitend, wenn der Schalter 2 geschlossen ist, so daß nur durch den Schalter 2 beide Relais eingeschaltet werden können.

Zur Vertiefung unserer Erkenntnisse wollen wir eine weitere Schaltung heranziehen, die als schematisierter Auszug aus einer in der Datenverarbeitung

¹⁾ Modulator = Gerät, mit dessen Hilfe elektrische Schwingungen im Rhythmus von Sprache oder Musik in der Stärke beeinflusst werden.

gebräuchlichen Schaltung anzusehen ist. Mit Hilfe der Schalter 1 bis 7 sollen drei Lampen derart geschaltet werden, daß die Summe der Zahlen an den jeweils eingeschalteten Lampen mit der Nummer des Schalters übereinstimmt.

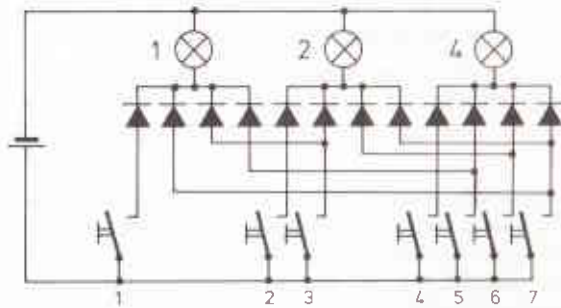


Abb. 3.32 — Codierschaltung zur Umsetzung von Dezimalzahlen in Dualzahlen

Zur Funktionsbeschreibung können wir eine einfache Wertetabelle aufstellen:

Nr. des geschlossenen Schalters	Nr. an den eingeschalteten Lampen	Summe der Lampen-Nr.
1	1	1
2	2	2
3	1 + 2	3
4	4	4
5	1 + 4	5
6	2 + 4	6
7	1 + 2 + 4	7

Diese Schaltung ermöglicht in einfacher Weise die Umsetzung von Dezimalzahlen ¹⁾ in Dualzahlen ²⁾. Sie wird auch Codierschaltung ³⁾ bezeichnet.

Ein wesentlicher Vorteil, der sich aus der Anwendung von Entkopplerdioden ergibt, ist die erhebliche Einsparung von Schalterkontakten. Wollte man die obige Codierschaltung ohne Entkopplerdioden — also nur mit Schalterkontakten — verwirklichen, so müßten die Schalter 3, 5, 6 und 7 bereits zwei bzw. drei voneinander isolierte Kontakte enthalten. In größeren

¹⁾ Dezimalzahlen = Zahlen des Zehnersystems, wie wir sie beim üblichen Rechnen und Zählen anwenden.

²⁾ Dualzahlen = Zahlen eines nur aus Zweierschritten bestehenden Zahlensystems, das in der elektronischen Datenverarbeitung angewendet wird.

³⁾ codieren = verschlüsseln; in eine andere Sprache umsetzen.

Schaltungen würde der Bedarf an Kontakten und damit auch der Raumbedarf zu hoch, um solche Schaltungen wirtschaftlich herstellen zu können.

In der folgenden Versuchsschaltung werden Dioden als Schalter und Entkoppler eingesetzt. Die Schaltung wird mit Wechselstrom betrieben.

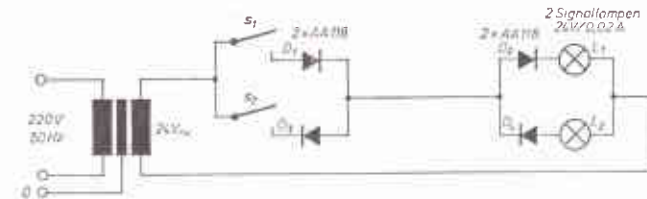


Abb. 3.33 — Dioden als Schalter und Entkoppler

Wären die Dioden D_1 und D_3 nicht vorhanden, so würden sowohl beim Betätigen des Schalters S_1 als auch des Schalters S_2 beide Lampen aufleuchten, und zwar abwechselnd für die Dauer je einer Halbwelle. Mit Hilfe der Dioden D_1 und D_3 werden aber die beiden Lampenstromkreise gegeneinander entkoppelt. Ist der Schalter S_1 geschlossen, so kann z. B. nur die positive Halbwelle des Wechselstroms über $D_1 - D_2 - L_1$ fließen; denn für die negative Halbwelle ist D_4 zwar durchlässig, D_1 jedoch gesperrt. Somit kann nur die Lampe L_1 aufleuchten. Wird nur der Schalter S_2 geschlossen, so sind die Dioden D_3 und D_4 für die negative Halbwelle des Wechselstroms durchlässig. Damit ist der Stromkreis für die Lampe L_2 geschlossen. Dagegen verhindert D_3 jetzt einen Stromfluß für die positive Halbwelle über D_2 und Lampe L_1 .

Auf diese Weise können über eine einzige Doppelleitung ohne Relais zwei Stromkreise unabhängig voneinander ein- und ausgeschaltet werden. Dieses auch als Halbwellensteuerung bezeichnete Verfahren wird u. a. auch bei elektronischen Modelleisenbahnen (zwei E-Loks auf einer Schiene) und Hausklingelanlagen (zwei Klingelanlagen an einer Klingelleitung) angewandt.

3.3.4. Die Diode als Spannungsbegrenzer

Ein PN-Übergang wird in Durchlaßrichtung erst leitend, wenn die angelegte Spannung größer als die Diffusionsspannung ist. Im Widerstandsverhalten äußert sich das so, daß der dynamische Widerstand einer Diode bis zum Kennlinienknick sehr groß (Größenordnung etwa 10 k Ω bis 100 k Ω), darüber jedoch verhältnismäßig klein (Größenordnung etwa 0,1 k Ω bis 1 Ω) ist. Diese Eigenschaft macht man sich zunutze, um Spannungsschwankungen zu begrenzen. Eine Diode wird dazu immer parallel zum Verbraucher geschaltet; ihre Wirkungsweise als Begrenzer zeigt folgender Versuch. Als

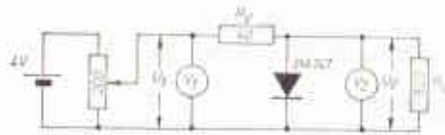


Abb. 3.34 — Diode als Spannungsbegrenzer

Wird die Versorgungsspannung erhöht, so steigt die Spannung an der Diode bis ca. 0,6 Volt im gleichen Umfang wie die Versorgungsspannung. Eine weitere Erhöhung der Versorgungsspannung U_1 hat jedoch nur noch einen äußerst geringen Anstieg der Diodenspannung U_D zur Folge.

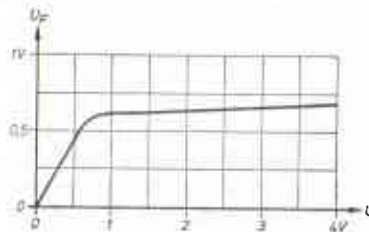


Abb. 3.35 — Regeldiagramm für eine Begrenzerdiode

Während also die Versorgungsspannung von ca. 0,65 bis 4 Volt steigt, ändert sich die Diodenspannung und damit die Spannung am Verbraucher kaum noch. Das anhand der Meßwerte gezeichnete Regeldiagramm bestätigt dieses Verhalten. Für die praktische Anwendung der Diode als Begrenzer einige Beispiele:

Beispiel 1: Gehörschutz

Die Hörkapsel eines Fernsprechapparats benötigt für eine ansehnlich lautstarke Wiedergabe eine Tonfrequenz-Wechselspannung von etwa 0,2 Volt effektiv. Der Spitzenwert einer sinusförmigen Wechselspannung von 0,2 V_{eff} beträgt rd. 0,3 Volt. Treten in einer Fernsprehverbindung Störspannungen auf, so können diese erheblich höhere Werte erreichen (Knackgeräusche). Das menschliche Ohr ist gegenüber plötzlichen, lautstarken Schallstößen sehr empfindlich. Jeder hat schon einmal die Erfahrung gemacht, daß die Hörempfindung nach einem lauten Knall geringer wird; man hat dann das Gefühl, für kurze Zeit taub zu sein. Solche Erscheinungen würden die Verständlichkeit beim Fernsprechen stark beeinträchtigen. Da es sich beim Fernsprechen um Wechselspannungs-Vorgänge handelt, werden parallel zur Hörkapsel zwei Dioden in entgegengesetzter Durchlaßrichtung als sogenannte **Gehörschutzgleichrichter** geschaltet; man spricht bei dieser Schaltung auch von Antiparallelschaltung.

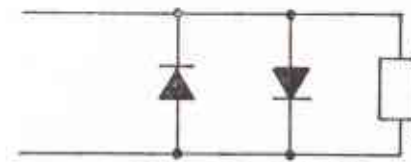


Abb. 3.36 — Gehörschutzgleichrichter

Für alle Spannungswerte, die die Diffusionsspannung überschreiten, stellen die Dioden einen niederohmigen Nebenschluß zur Hörkapsel dar. Als Gehörschutzgleichrichter eignen sich sehr gut Dioden auf Selenbasis, da ihre Diffusionsspannung etwa 0,4 Volt beträgt. Die zur Ansteuerung der Hörkapsel notwendige Spannung von 0,3 V_{max} wird von den Selendiode noch nicht beeinflußt. Da eine Begrenzerschaltung grundsätzlich nur wirksam ist, wenn ein Vorwiderstand vorhanden ist, fällt der Spannungsüberschuß an den in der Schaltung liegenden Reihenwiderständen (z.B. R_1 der Induktionsspule) ab.

Beispiel 2: Spannungsstabilisierung

Transistorverstärker werden vielfach aus Batterien gespeist; die Klemmenspannung einer Batterie sinkt jedoch mit zunehmender Entladung. Das hat zwangsläufig eine Verminderung der Verstärkerwirkung zur Folge. Da — wie wir noch genauer untersuchen werden — die richtige Arbeitsweise eines Transistors sehr wesentlich von der Höhe der Basis-Emitter-Spannung abhängt, kann man diese mit Hilfe von Begrenzerdioden stabilisieren. Diese Schaltungsmaßnahme wird überwiegend in der Endstufe eines Transistorverstärkers vorgenommen. Nebenstehender Schaltungsausschnitt zeigt eine Gegentaktendstufe mit sogenannten Komplementärtransistoren. Die beiden in Reihe geschalteten Basis-Emitter-Strecken der Transistoren liegen parallel zur Begrenzerdiode und erhalten eine vom Entladezustand der Batterie nahezu unabhängige Spannung. Da die Batteriespannung bis zur zulässigen Entladegrenze immer noch wesentlich höher als die Diffusionsspannung der Diode ist, wirkt sich die Stabilisierungsmaßnahme bis zur Entladegrenze der Batterie aus. Muß die stabilisierte Spannung etwas größer sein als die Diffusionsspannung, so können auch zwei oder mehrere Dioden in Reihe geschaltet werden, ist die benötigte Spannung dagegen kleiner, so kann man einen Spannungsteiler nachschalten.

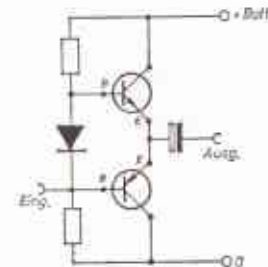


Abb. 3.37 — Stabilisierung der Basis-Emitter-Spannung einer Transistorstufe

Beispiel 3: Impulsformung

Die Begrenzereigenschaft kann auch zur **Impulsformung** ausgenutzt werden. Von dieser Möglichkeit wird in der Elektronik sehr häufig Gebrauch gemacht. Eine einfache Versuchsschaltung für einen Impulsformer mit Begrenzerdiode zeigt Abb. 3.38. Der vom Transformator abgegebene Wechselstrom wird durch die Diode AA 118 gleichgerichtet (Einweggleichrichtung). Bei geöffnetem Schalter s zeigt sich auf dem Bild-

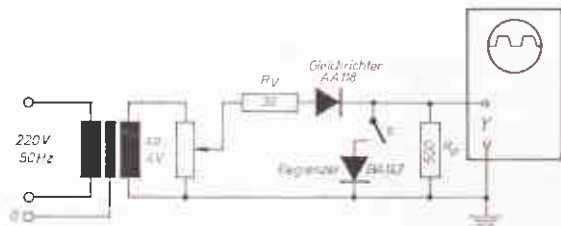


Abb. 3.38 — Begrenzerdiode als Impulsformer

schirm des Oszilloskops das bekannte Diagramm der pulsierenden Gleichspannung an R_n nach Einweggleichrichtung. Wird der Schalter s geschlossen, so beschneidet die Begrenzerdiode BA 147 die über ca. 0,6 Volt hinausgehenden Spannungsspitzen. Die

Begrenzerdiode stellt für Spannungen oberhalb ihrer Diffusionsspannung einen niederohmigen Nebenschluß zum Verbraucher R_n dar. Das Oszillogramm zeigt Sinuspulse, deren Spitzen „abgeschnitten“ sind. Auf diese Weise sind also die Sinuspulse zu trapezförmigen Impulsen umgeformt worden. Die Impulsformung wird deutlicher, je höher die Versorgungsspannung mit dem Potentiometer eingestellt wird (zulässige Strombelastung der Dioden beachten!).



Abb. 3.39 — Oszillogramm der Impulse

3.3.5. Die Diode als Leistungsbegrenzer

Wird eine Gleichrichter-Diode in einen Wechselstromkreis eingeschaltet, so läßt sie wegen ihrer Stromrichtungsabhängigkeit nur je eine Halbwelle einer Wechselstromperiode durch. Wir wollen einmal annehmen, daß die Diode nur für die positiven Halbwellen durchlässig ist. Liegt in Reihe mit der Diode ein Verbraucher, so wird dieser nur von den positiven Halbwellen des Wechselstroms durchflossen; während der negativen Halbwellen ist er stromlos. Daraus ergibt sich nun, daß der Verbraucher bei gleicher Wechselspannung gegenüber der Schaltung ohne Diode nur noch die halbe Leistung aufnimmt.

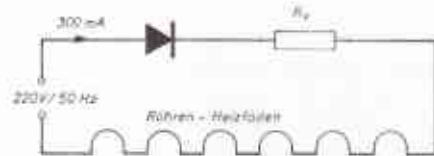


Abb. 3.40 — Diode als Leistungsbegrenzer

Die Diode hat hier also die Wirkungsweise eines Leistungsbegrenzers. Ihr Vorteil gegenüber anderen Arten der Leistungsbegrenzung in Wechselstromkreisen (z.B. durch Vorwiderstände, Drosseln oder Kondensatoren) besteht in der geringen Baugröße moderner Gleichrichter-Dioden. Sehr wesentlich ist aber auch, daß die Diode als Leistungsbegrenzer selbst so gut wie keine Leistung verbraucht. Die Diode ist im Durchlaßbereich sehr niederohmig und im Sperrbereich sehr hochohmig. Der einzige Nachteil ist, daß sie die Leistungsaufnahme des Verbrauchers grundsätzlich halbiert; es ist daher nicht ohne weiteres möglich, Zwischengrößen einzustellen. Die Diode hat sich inzwischen als Wechselstrom-Leistungsbegrenzer einen festen Platz erobert; dazu zwei Beispiele:

Beispiel 1: Heizleistungsbegrenzung in Fernsehgeräten

Fernsehgeräte sind heute teilweise mit Transistoren und teilweise mit Röhren bestückt. Die Heizfäden aller Röhren werden dabei in Reihe geschaltet. Somit ergibt sich die Gesamtheizspannung aus der Summe der Heizspannungen der einzelnen Röhren. Wegen der immer geringer werdenden Zahl von Röhren in einem Fernsehgerät ist aber auch der Heizspannungsbedarf geringer. Da ohne Transformator gearbeitet wird, muß der Spannungsunterschied zwischen Netzspannung und Gesamtheizspannung an Vorwiderständen abfallen. Diese Vorwiderstände nehmen aber selbst elektrische Leistung auf und setzen sie in Wärme um. Die Erwärmung ist — abgesehen vom auftretenden Leistungsverlust — sehr unerwünscht, da die Wirkungsweise anderer Bauteile (vor allem der Transistoren) dadurch stark beeinträchtigt wird.



**Abb. 3.41 — Heizleistungsbegrenzung
in einem Fernsehgerät**

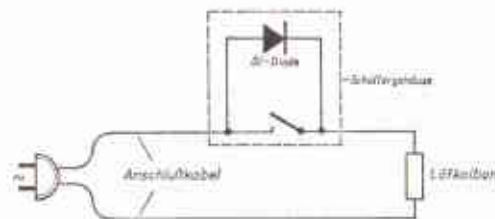
Damit wird der Energieverbrauch und die Erwärmung der Widerstände aufgrund der Beziehung $I^2 \cdot R$ ebenfalls geringer; eine entsprechende Schaltung ist in Abb. 3.41 wiedergegeben.

Beispiel 2: Leistungsrosselung bei LötKolben

Wird ein elektrischer LötKolben längere Zeit betrieben, erwärmt sich die Kupferspitze so stark, daß sie verzündet und sich dann nicht mehr einwandfrei verzinnt läßt, wodurch Lötarbeiten stark erschwert werden. Die im Laufe längerer Betriebszeit zu stark angestiegene Temperatur der Kupferspitze läßt aber auch das geschmolzene Lötzinn sehr schnell an der Oberfläche oxydieren; beim Löten gelangt das Oxid leicht an die Verbindungsstelle. Die Folge sind fehlerhafte Verbindungen, die in krassen Fällen sogar zur Unterbrechung führen. Um das zu vermeiden, ist es natürlich möglich, den LötKolben zeitweise auszuschalten, damit er abkühlen kann. Das führt jedoch zu unliebsamen Unterbrechungen der Lötarbeiten. Hier kommt uns die Leistungsbegrenzer-Eigenschaft einer Gleichrichter-Diode zu Hilfe.

Schalten wir in den LötPausen vor den LötKolben eine Gleichrichter-Diode, so wird die Leistungsaufnahme des LötKolbens auf die Hälfte reduziert. Der LötKolben bleibt damit zwar warm, er wird jedoch nicht überhitzt. Schon kurz nach Überbrückung der Diode hat er seine volle Betriebstemperatur wieder erreicht.

Die praktische Ausführung dieser Schaltungsmaßnahme ist sehr einfach: Man baut in das Anschlußkabel des LötKolbens einen Schnurschalter ein, der auch für Wohnzimmerleuchten handelsüblich ist. In diesen Schalter wird parallel zu den Schalterkontakten eine Silizium-Leistungsdioden eingesetzt, wie es die folgende Schaltung zeigt.



**Abb. 3.42 — Leistungsbegrenzung
bei einem LötKolben**

Durch Vorschalten einer Diode wird die Leistungsaufnahme der Röhrenheizungen ohne wesentlichen Leistungsverlust auf die Hälfte herabgesetzt. Ist der Heizleistungsbedarf noch geringer, so werden zusätzlich Vorwiderstände verwendet; sie haben aber einen wesentlich geringeren Widerstand, als wenn sie ausschließlich die Heizleistung begrenzen müßten.

Wird bei gestecktem Anschlußstecker der Schalter unterbrochen, so wird automatisch die Diode in den Stromkreis geschaltet und die Leistungsaufnahme des LötKolbens halbiert. Der Raumbedarf der modernen Silizium-Leistungsdioden ist so gering, daß sie sich leicht im Schaltergehäuse unterbringen lassen. Die Diode muß mit dem Betriebsstrom des LötKolbens belastbar sein. Dabei ist jedoch zu berücksichtigen, daß der

Widerstand des LötKolbens im kalten Zustand geringer ist als im Betriebszustand. Wir rechnen daher sicherheitshalber mit dem Faktor 1,5. So ergibt sich z.B. für einen LötKolben 220 V / 30 W ein Betriebsstrom von

$$I = \frac{P}{U} = \frac{30 \text{ W}}{220 \text{ V}} = 0,137 \text{ A}$$

und bei 1,5facher Sicherheit von $1,5 \cdot 0,137 \text{ A} = 0,205 \text{ A}$.

Selbstverständlich muß die Diode für eine Sperrspannung von 220 V_{eff} ausgelegt sein. Noch eleganter könnte das Problem durch einen Gabelschalter gelöst werden, der durch das Ablegen des LötKolbens betätigt wird. Wegen der mit dem Umgang mit elektrischer Installation verbundenen Gefahren wird jedoch vom Selbstbau solcher Gabelschalter dringend abgeraten.

3.3.6. Die Diode als Gegenzeile

Der Strombedarf einer Fernmeldeanlage schwankt im Verlauf eines Tages sehr stark; damit treten zwangsläufig in der Stromversorgungsanlage Spannungsschwankungen auf. Zum sicheren Betrieb von Fernmeldeanlagen ist jedoch die Einhaltung bestimmter Grenzen erforderlich. In Stromversorgungsanlagen kann die Spannung entweder im Netzgerät oder am Verbraucher geregelt werden. Heute ist vielfach die lastabhängige Zu- oder Abschaltung von Gegenzellen gebräuchlich. Als Gegenzellen wurden früher Bleisammler oder alkalische Gegenzellen eingesetzt, sie sind jedoch längst durch Halbleiter-Gleichrichter ersetzt worden. Die Gegenzellen werden in Gleichstromversorgungsanlagen in Durchlaßrichtung zwischen Stromversorgungsanlage und Verbraucher geschaltet; sie können durch Schalter überbrückt werden. Die Wirkungsweise des Halbleitergleichrichters als „Gegenzeile“ beruht auf dem Kennlinienknick im Durchlaßbereich, hervorgerufen durch die Diffusionsspannung. Schaltet man einen Halbleitergleichrichter — z.B. einen Silizium-Gleichrichter — in Durchlaßrichtung zwischen Gleichstromquelle und Verbraucher, so ist die Spannung am Verbraucher um den Betrag der Diffusionsspannung geringer. Die Diffusionsspannung ist ja — wie bereits bekannt — der angelegten Spannung entgegengerichtet. Je nach Höhe der erforderlichen Spannungsänderung werden eine oder mehrere Gegenzellen durch Schalter in den Verbraucherstromkreis geschaltet. Da die Trockengleichrichter auch im Durchlaßbereich einen — wenn auch kleinen — inneren Widerstand besitzen, ist der wirksame Spannungsverlust geringfügig stromabhängig. Er schwankt etwa bei Selen-Gegenzellen zwischen 0,4 und 0,7 V, bei Silizium-Gegenzellen zwischen 0,6 und 0,8 V je nach Strombelastung.

3.4. Referenzdioden (Z-Dioden)

3.4.1. Grundsätzlicher Aufbau und Wirkungsweise

Die bisher untersuchten Dioden werden im Sperrzustand sofort zerstört, wenn die zulässige Sperrspannung um einen geringen Betrag überschritten wird. Infolge der im Sperrzustand an der Sperrzone liegenden hohen elektrischen Feldstärke werden dann plötzlich Elektronen aus ihren festen Bindungen gerissen (**Zener effekt**). Jeder der so frei gewordenen Elektronen kann unter bestimmten Bedingungen durch das elektrische Feld so stark beschleunigt werden, daß es beim Auftreffen auf andere Atome dort mehrere Elektronen herausschlägt. Die Zahl der freien Elektronen nimmt daher lawinenartig zu. Diesen Prozeß des lawinenartigen Durchbruchs nennt man **Avalanche-Effekt**. Seine Entdeckung wurde irrtümlicherweise dem amerikanischen Physiker Zener zugeschrieben, weshalb Dioden, die diesen Effekt ausgeprägt zeigen, Zenerdioden genannt wurden. Heute setzen sich die Bezeichnungen Referenzdioden oder Z-Diode mehr und mehr durch. Da-



Abb. 3.43

neben werden allerdings charakteristische Größen und Eigenschaften nach wie vor nach Zener benannt, z.B. Zenerspannung und Zenerwiderstand. Als Schaltzeichen wird für Referenzdioden nebenstehendes Symbol verwendet.

Eine Referenzdiode ist äußerlich wie eine Universaldiode aufgebaut und enthält einen N-Siliziumkristall, dem Aluminium (3wertig) einlegiert ist. Der dabei entstandene PN-Übergang zeigt gegenüber herkömmlichen PN-Übergängen besondere Eigenschaften. Wird die angelegte Spannung einer Referenzdiode im Sperrbereich über die Durchbruchspannung hinaus erhöht, so tritt keine Zerstörung ein, solange eine bestimmte zulässige Verlustleistung nicht überschritten wird. Ein weiterer Unterschied zu Universaldioden besteht darin, daß der Stromanstieg beim Überschreiten der Durchbruchspannung sehr stark ist, der Widerstand also plötzlich sehr klein wird. Der **Bereich des Stromstellanstiegs wird Zenerbereich**, die Spannung, bei der der Durchbruch erfolgt, **Zenerspannung** genannt. Es ist im allgemeinen nur ein kleiner Übergangsbereich zwischen gesperrtem Zustand und dem Zenerbereich vorhanden. Wir wollen das Verhalten einer Referenzdiode an ihrer Kennlinie veranschaulichen. Die Kennlinie kann übrigens mit der gleichen Meßschaltung, wie für den Sperrbereich einer Universaldiode, aufgenommen werden (vgl. Abb. 3.13). Sie ist in Abb. 3.44 für eine Referenzdiode mit einer Zenerspannung von $U_Z = 5\text{ V}$ und einer zulässigen Verlustleistung von $P_V = 100\text{ mW}$ dargestellt und läßt erkennen:

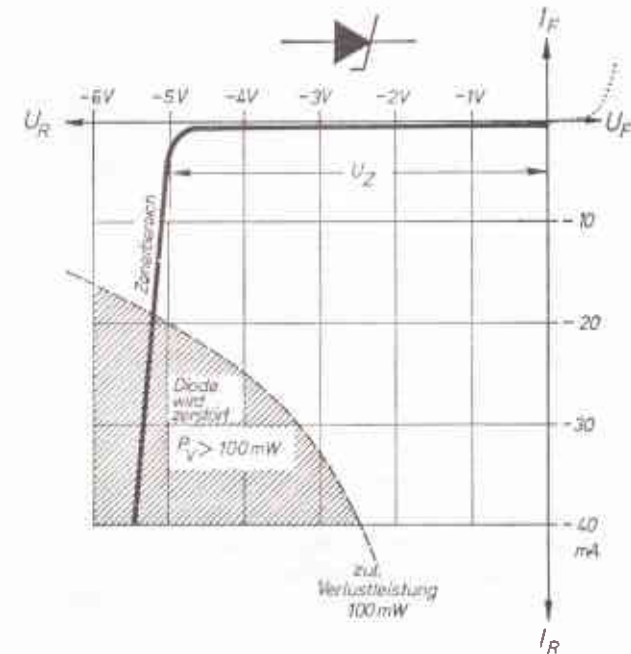


Abb. 3.44 — Kennlinie einer Referenzdiode (Z-Diode)

- Bis zur Höhe der Zenerspannung fließt ein sehr geringer Sperrstrom. Die Referenzdiode ist noch sehr hochohmig bis in die Größenordnung von einigen Megohm.
- In der Nähe der Zenerspannung beginnt der Stromanstieg zunächst nur sehr langsam (Übergangsbereich).
- Beim Erreichen bzw. Überschreiten der Zenerspannung steigt der Strom stark an. Die große Steilheit der Kennlinie im Zenerbereich läßt auf einen sehr kleinen dynamischen Widerstand schließen. Dieser Widerstand wird auch als Zenerwiderstand r_z bezeichnet und kann bis auf etwa 1 Ohm absinken.

Referenzdioden können für Zenerspannungen zwischen etwa 3 V und 100 V hergestellt werden. Zur Lösung von Stabilisierungs- und Begrenzeraufgaben unterhalb 3 V verwendet man normale Dioden im Durchlaßbereich (Begrenzerdioden, wie bereits beschrieben), für größere Spannungswerte werden je nach Bedarf mehrere Referenzdioden in Reihe geschaltet.

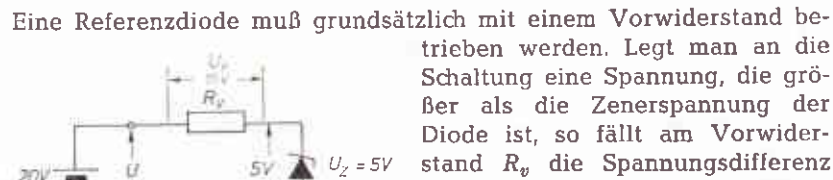


Abb. 3.45 — Grundsätzliche Schaltung einer Referenzdiode

die Referenzdiode im dargestellten Schaltungsbeispiel überlastet und zerstört. Das Verhalten wird noch an Versuchsschaltungen nachgewiesen.

Referenzdioden werden grundsätzlich in Sperrichtung betrieben. Sie dürfen im Gegensatz zu anderen Dioden unter Beachtung der zulässigen Verlustleistung auch an Spannungen oberhalb der Durchbruchspannung liegen. Ihr dynamischer Widerstand ist im Zenerbereich sehr klein. Zum Betrieb einer Referenzdiode gehört immer ein Vorwiderstand.

Folgende Größen kennzeichnen das Verhalten einer Referenzdiode:

- Die Zenerspannung U_z
(wird für einen bestimmten Wert des Zenerstroms, z.B. 5 mA, 25 mA oder 100 mA angegeben),
- der Zenerwiderstand r_z
(= dynamischer Widerstand im Zenerbereich, meistens bei $I_z = 5$ mA oder 100 mA gemessen),
- die zulässige Verlustleistung P_v
(= Produkt aus anliegender Spannung und fließendem Strom),
- die zulässige Höchsttemperatur
(sie beträgt fast ausschließlich 150°C).

Die spannungsstabilisierende Wirkung einer Referenzdiode ist, wie wir noch untersuchen werden, um so besser, je kleiner ihr dynamischer Widerstand r_z im Zenerbereich ist. Ermittelt man die Zener-

trieben werden. Legt man an die Schaltung eine Spannung, die größer als die Zenerspannung der Diode ist, so fällt am Vorwiderstand R_v die Spannungsdifferenz

$$U - U_z = U_v \quad \text{ab;}$$

denn für Spannungen oberhalb der Zenerspannung ist die Referenzdiode sehr niederohmig. Beim Fehlen des Vorwiderstands würde

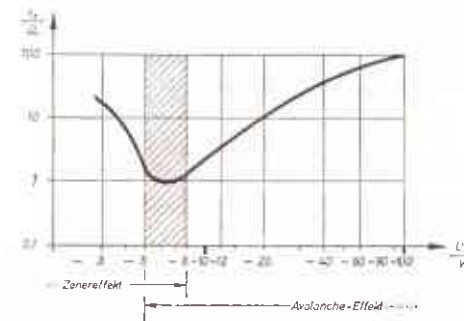


Abb. 3.46 — Abhängigkeit des Zenerwiderstands von der Zenerspannung

bruch überwiegend als Folge des Zenereffekts auftritt, während bei Dioden mit Zenerspannungen oberhalb von 5 V hauptsächlich der Avalanche-Effekt für den plötzlichen Stromanstieg verantwortlich ist. Man erkennt, daß im Bereich 5 bis 8 V beide Effekte nebeneinander auftreten können. Darum ergibt sich hier beim Erreichen der Zenerspannung ein besonders steiler Stromanstieg.

Wegen ihres besonders kleinen Zenerwiderstands werden für Stabilisierungsschaltungen bevorzugt Referenzdioden mit Zenerspannungen zwischen 5 und 8 V verwendet. Für höhere Spannungen schaltet man vielfach Referenzdioden in Reihe.

Beispiel:

Der Zenerwiderstand der Referenzdioden aus der Typenserie BZY 92 beträgt für $U_z = 6\text{ V}$ $r_z = 1\ \Omega$, dagegen für $U_z = 12\text{ V}$ $r_z = 4\ \Omega$. Soll eine Spannung von 12 V stabilisiert werden, so könnte entweder eine 12-V-Diode oder die Reihenschaltung zweier 6-V-Dioden eingesetzt werden. Dabei ergeben sich folgende Zenerwiderstände:

- für die 12-V-Diode $r_z = 4\ \Omega$ und
- für zwei in Reihe geschaltete 6-V-Dioden $r_z = 2\ \Omega$.

Trotz der Reihenschaltung ist der Widerstand bei Verwendung zweier 6-V-Dioden nur halb so groß wie der einer einzelnen 12-V-Diode.

3.4.2. Spannungs-Stabilisierungsschaltungen mit Referenzdioden

Der Einfachheit halber wollen wir die Referenzdiode künftig **Z-Diode** nennen. Ihre verbreitetste Anwendung findet die Z-Diode zur Stabilisierung von Spannungen. Für die folgenden Bemessungsbeispiele

wird die Z-Diode BZY 92 C 6V2 zugrunde gelegt, über die das Datenblatt folgende Angaben enthält:

$$\text{BZY 92 C 6V2} \left\{ \begin{array}{l} U_z = 6,2 \text{ V} \pm 5 \% \\ P_v = 1 \text{ W} \\ r_z = 1 \ \Omega \\ \vartheta_{zu} = 150 \text{ } ^\circ\text{C} \end{array} \right.$$

Um die spannungsstabilisierende Wirkung einer Z-Diode besser zu verstehen, wollen wir ihre Kennlinie etwas genauer untersuchen. Die

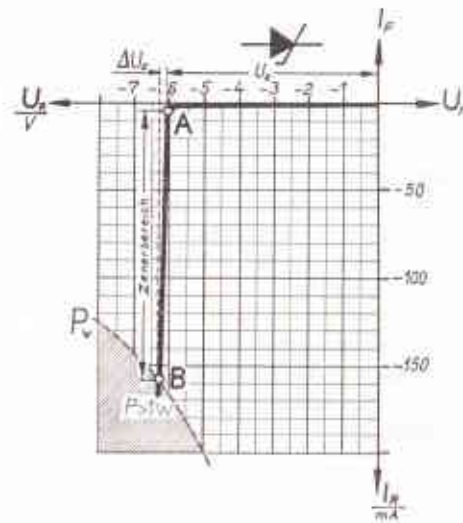


Abb. 3.47 — Kennlinie der Z-Diode
BZY 92 C 6V2

(Sperr-)Kennlinie einer Z-Diode zeigt, daß die an ihr liegende Spannung als nahezu konstant anzusehen ist, solange die Z-Diode im Zenerbereich betrieben wird. Damit sie jedoch nicht zerstört wird, darf die Diode nur bis zur zulässigen Leistungsgrenze (Leistungshyperbel) belastet werden. Ohne Gefahr für die Z-Diode bleibt die an ihr liegende Spannung nahezu gleich, wenn wir dafür sorgen, daß der Arbeitspunkt zwischen den Punkten A und B auf der Kennlinie liegt. Bei der Bemessung von Stabilisierungsschaltungen muß demnach von zwei Grenzfällen ausgegangen werden:

Grenzfall A: Die Z-Diode nimmt praktisch keinen Strom auf, liegt aber an einer Spannung in Höhe von U_z .

Grenzfall B: Die Spannung an der Z-Diode liegt nur geringfügig höher als im Arbeitspunkt A. Jedoch nimmt die Z-Diode hier den maximal zulässigen Strom auf.

Während also die Stromaufnahme von nahezu Null bis $I_{z\max}$ steigt, ändert sich die Spannung nur um einige Zehntel Volt. An der darge-

stellten Kennlinie ist auch leicht zu erkennen, daß die Spannungsänderung im Zenerbereich um so kleiner ist, je steiler die Kennlinie verläuft (je kleiner also der Zenerwiderstand r_z ist).

Die Zenerspannung U_z der Z-Diode BZY 92 C 6V2 wird in den folgenden Schaltungs- und Berechnungsbeispielen der Einfachheit halber mit 6 V angenommen.

3.4.2.1. Die Z-Diode an veränderlicher Spannung ohne Verbraucher

Da die Wirkungsweise einer Spannungs-Stabilisierungsschaltung am einfachsten zu übersehen ist, wenn kein Verbraucher angeschlossen ist, soll zunächst eine einfache Schaltung untersucht werden. Sie könnte als Konstantspannungsquelle bezeichnet werden, wobei davon ausgegangen wird, daß die Eingangsspannung in weiten Grenzen schwankt.

Eine Trockenbatterie gilt als entladen, wenn ihre Spannung auf den halben Wert der Nennspannung abgesunken ist. Bei einer batteriebetriebenen Schaltung müssen also Spannungsänderungen im Verhältnis 2 : 1 in Kauf genommen werden, wenn die Batterie bis zur zulässigen Entladegrenze ausgenutzt werden soll. Für den einwandfreien Betrieb eines Transistorverstärkers sind jedoch solche Spannungsunterschiede zu groß. Je weiter die Versorgungsspannung abnimmt, desto stärker werden die im Verstärker auftretenden Verzerrungen. Mit Z-Dioden-Schaltungen gelingt es, selbst große Spannungsänderungen von Batterien oder anderen Stromquellen auszugleichen.

Wir haben bereits festgestellt, daß die Spannung an einer Z-Diode konstant bleibt, solange der Arbeitspunkt zwischen den Grenzfällen A und B liegt. Mit der gewählten Z-Diode BZY 92 C 6V2 liegt die untere Spannungsgrenze mit rd. 6 V fest. Um eine Batterie voll ausnutzen zu können, müßte ihre Nennspannung im frischen Zustand 12 V betragen. Im Versuchsaufbau wird zur Spannungsregelung ein Potentiometer benutzt.

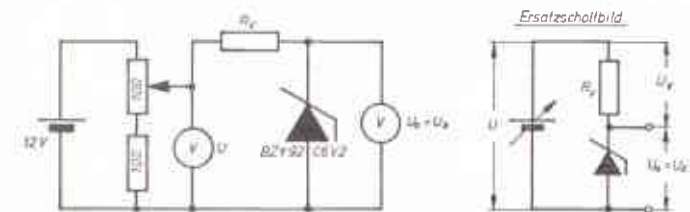


Abb. 3.48 — Spannungs-Stabilisierung mit Hilfe einer Z-Diode

Damit die Z-Diode nicht überlastet werden kann, muß der Vorwiderstand R_V für den Grenzfall B ($I_{Z\max}$) bemessen werden. Bei voller Batteriespannung $U = 12\text{ V}$ muß am Vorwiderstand der Spannungsüberschuß

$$U_V = U - U_Z \quad \text{abfallen.}$$

Dabei darf der Strom nicht größer als $I_{Z\max}$ sein.

Also:
$$R_V = \frac{U_V}{I_{Z\max}}$$

$$U_V = 12\text{ V} - 6\text{ V} = 6\text{ V}$$

$$I_{Z\max} = \frac{P_V}{U_Z} = \frac{1\text{ W}}{6\text{ V}} = 0,167\text{ A}$$

Aus Sicherheitsgründen rechnen wir mit $I_{Z\max} = 0,15\text{ A}$.

$$R_V = \frac{6\text{ V}}{0,15\text{ A}} = 40\ \Omega.$$

Da die Z-Diode im Grenzfall A praktisch stromlos wird, tritt hier auch an R_V kein Spannungsabfall mehr auf. Wird die Eingangsspannung der Stabilisierungsschaltung stufenweise von 12 V auf 6 V gesenkt, so läßt sich anhand der Meßwertreihe folgendes Regeldiagramm darstellen:

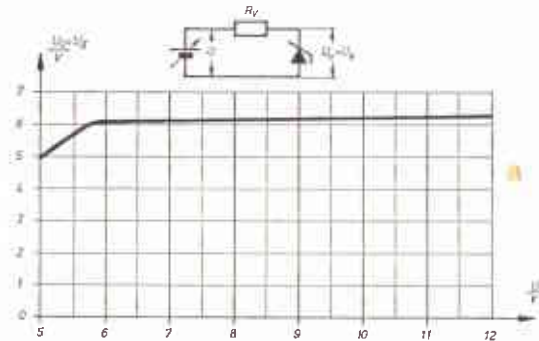


Abb. 3.49 — Regeldiagramm der Spannungs-Stabilisierungsschaltung

Ergebnis: Trotz Änderung der Eingangsspannung von 12 V auf 6 V (also um 50 %) ist die Ausgangsspannung der Schaltung praktisch konstant 6 V geblieben.

3.4.2.2. Die Z-Diode an veränderlicher Spannung mit angeschlossenem Verbraucher

Für das folgende Bemessungsbeispiel wird wieder davon ausgegangen, daß die Versorgungsspannung in weiten Grenzen schwankt. Am Ausgang der Stabilisierungsschaltung liegt ein Verbraucher. In der Versuchsschaltung ist es eine Glühlampe 6 V/0,1 A. Wie im ersten Beispiel soll die höchste Spannung 12 V betragen. Zur meßtechnischen Untersuchung wird die Eingangsspannung mit einem Potentiometer geregelt.

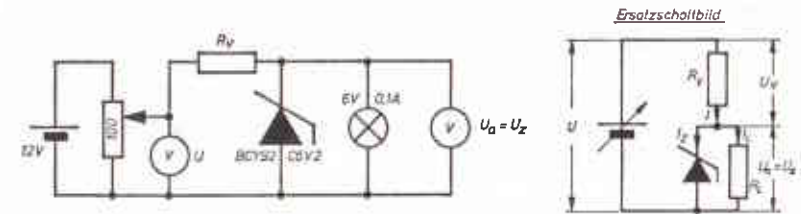


Abb. 3.50 — Stabilisierung der Versorgungsspannung für einen Verbraucher

Die Z-Diode nimmt den größten Strom auf, wenn die Eingangsspannung ihren Höchstwert (12 V) hat. Damit sie nicht zerstört wird, muß der Vorwiderstand R_V für den Grenzfall B ($I_{Z\max}$) bemessen werden.

Grenzfall B: ($I_{Z\max}$)
$$R_V = \frac{U_{V\max}}{I_{\max}}$$

$$U_{V\max} = U - U_Z$$

$$U_{V\max} = 12\text{ V} - 6\text{ V} = 6\text{ V}$$

$$I_{\max} = I_{Z\max} + I_L$$

$$I_{\max} = 0,15\text{ A} + 0,1\text{ A} = 0,25\text{ A}$$

$$R_V = \frac{6\text{ V}}{0,25\text{ A}} = 24\ \Omega.$$

Der Vorwiderstand muß einen Wert von 24 Ω haben.

In dieser Schaltung fließt auch Strom, wenn die Eingangsspannung unter die Zenerspannung U_Z der Z-Diode absinkt, und zwar der Verbraucherstrom I_L . Soll die Verbraucherspannung konstant bleiben, so darf die Eingangsspannung nicht bis auf 6 V sinken. Da der Verbraucherstrom an R_V einen Spannungsabfall hervorruft, muß die Eingangsspannung mindestens um diesen Spannungsabfall größer sein. Die untere Grenze der Eingangsspannung kann für den Grenzfall A bestimmt werden, bei dem die Z-Diode zwar stromlos ist, der Verbraucherstrom aber nach wie vor fließt.

Grenzfall A:
($I_z = 0$)

$$U_{\min} = U_z + U_{V\min}$$

$$U_{V\min} = I_L \cdot R_V$$

$$U_{V\min} = 0,1 \text{ A} \cdot 24 \Omega = 2,4 \text{ V}$$

$$U_{\min} = 6 \text{ V} + 2,4 \text{ V} = \underline{8,4 \text{ V}}$$

Damit der Verbraucher grundsätzlich die Nennspannung 6 V erhält, darf die Eingangsspannung der Stabilisierungsschaltung nicht unter 8,4 V sinken.

Ergebnis: Die Verbraucherspannung bleibt nahezu konstant 6 V, wenn die Eingangsspannung von 12 V auf 8,4 V — also um 30 % fällt.

Dieses Bemessungsbeispiel läßt bereits die Tendenz erkennen, daß der Regelbereich einer Spannungs-Stabilisierungsschaltung mit Z-Diode bei Belastung kleiner wird.

3.4.2.3. Die Z-Diode an konstanter Eingangsspannung mit veränderlicher Stromaufnahme des Verbrauchers

Der Widerstand vieler Verbraucherschaltungen ändert sich im Betriebszustand in weiten Grenzen. So hängt beispielsweise der Widerstand und damit die Stromaufnahme einer Gegentakstverstärkerstufe mit Transistoren sehr stark von der Lautstärkeeinstellung ab. Da jede Stromquelle einen Innenwiderstand besitzt, treten infolge der schwankenden Stromentnahme auch Schwankungen der Klemmenspannung auf:

$$U_{kl} = U_0 - I \cdot R_i$$

Je größer der Verbraucherstrom I bzw. der Innenwiderstand R_i der Stromquelle ist, desto kleiner wird die Klemmenspannung U_{kl} .

Eine Schwankung der Versorgungsspannung führt häufig zu unliebsamen Störungen im Betrieb eines Verbrauchers. In Transistorverstärkern treten bei schwankender Versorgungsspannung Verzerrungen auf, deren Ursache oft irrtümlich im Verstärker selbst gesucht wird. Bei der Abstimmung von Resonanzkreisen mit Kapazitätsdioden würden infolge Spannungsschwankungen Schwankungen in der Resonanzfrequenz auftreten.

Auch hier kann uns die Z-Diode gute Dienste leisten. Für das folgende Bemessungsbeispiel gilt, daß die Leerlaufspannung U_0 und der Innenwiderstand R_i der Stromversorgung konstant bleiben, während sich der Verbraucherwiderstand ändert. Da Glühlampen am einfachsten an ihrer Helligkeit erkennen lassen, ob sich die anliegende Spannung ändert, verwenden wir für unsere Versuchsschaltung Glühlampen 6 V/0,05 A.

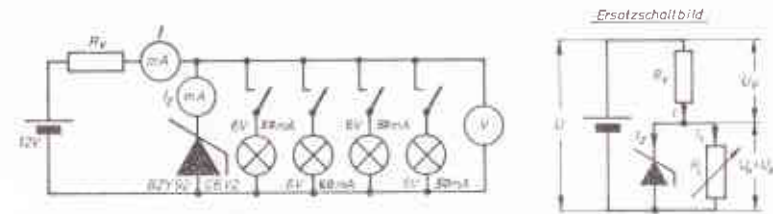


Abb. 3.51 — Spannungs-Stabilisierung mit Z-Diode bei veränderlichem Verbraucherwiderstand

In dieser Schaltung nimmt die Z-Diode den größten Strom auf, wenn der Verbraucherwiderstand am größten, bzw. wenn kein Verbraucher angeschlossen ist. Der Vorwiderstand R_V muß also für diesen Grenzfall bemessen werden, damit die Z-Diode nicht überlastet wird.

Grenzfall B:
($I_{z\max}$)
($R_L = \infty$)

$$R_V = \frac{U_V}{I_{z\max}}$$

$$U_V = U_0 - U_z$$

$$U_V = 12 \text{ V} - 6 \text{ V} = 6 \text{ V}$$

$$R_V = \frac{6 \text{ V}}{0,15 \text{ A}} = \underline{40 \Omega}$$

Ist der Innenwiderstand der Stromversorgung nicht vernachlässigbar klein, so kann der Wert des zusätzlichen Vorwiderstands um R_i verringert werden.

Dann gilt also

$$R_V = R_V - R_i$$

Werden jetzt nacheinander Glühlampen als Verbraucher eingeschaltet, so erhöht sich durch die Zunahme des Verbraucherstroms der Spannungsabfall an R_V . Die Teilspannung an der Parallelschaltung Z-Diode — Verbraucher wird somit kleiner. Der Arbeitspunkt der Z-Diode wandert auf der Kennlinie von Punkt B zum Punkt A. Soll die Verbraucherspannung konstant bleiben, so darf der Grenzfall A nicht unterschritten werden. Der Verbraucherstrom, d.h. die Zahl der eingeschalteten Glühlampen, darf also nicht beliebig erhöht werden.

Wieviel Glühlampen dürfen nun eingeschaltet werden, ohne daß die Spannung merklich absinkt? Da im Arbeitspunkt A bzw. Grenzfall A die Z-Diode selbst praktisch keinen Strom aufnimmt, wird der Spannungsabfall an R_V jetzt ausschließlich durch den Verbraucherstrom verursacht.

Grenzfall A: Der Spannungsabfall U_V am Vorwiderstand darf nicht größer als 6 V sein. Da der Widerstandswert für R_V mit 40 Ω festliegt, kann der höchstzulässige Verbraucherstrom berechnet werden:

$$I_{Lmax} = \frac{U_{Vmax}}{R_V} = \frac{6 \text{ V}}{40 \Omega} = \underline{0,15 \text{ A}}$$

Jede Lampe nimmt einen Strom von 0,05 A auf. Somit dürfen

$$n = \frac{I_{Lmax}}{I_{L1}} = \frac{0,15 \text{ A}}{0,05 \text{ A}} = \underline{3 \text{ Lampen}}$$

eingeschaltet werden.

Ergebnis: Die Ausgangsspannung der Stabilisierungsschaltung bleibt nahezu konstant, wenn der Verbraucherstrom zwischen 0 und 0,15 A schwankt, also $I_{Lmax} = I_{Lmax}$ ist.

Sobald in der Versuchsschaltung die vierte Lampe eingeschaltet wird, sinkt die Verbraucherspannung unter 6 V, was an der geringeren Helligkeit der Lampen bzw. an einem Spannungsmesser erkennbar wird. Interessant ist in der Versuchsschaltung die Messung der Ströme I und I_z . Während bis zu drei Lampen zu- oder abgeschaltet werden, bleibt der Gesamtstrom I praktisch konstant. Dagegen steigt der Zenerstrom I_z jeweils um etwa 50 mA, wenn eine Lampe ausgeschaltet wird und fällt um den gleichen Betrag, wenn eine Lampe eingeschaltet wird. Solange also der höchstzulässige Verbraucherstrom nicht überschritten wird, übernimmt die Z-Diode jeweils die Stromdifferenz zwischen dem höchsten und dem tatsächlichen Verbraucherstrom. Somit bleibt der Gesamtstrom I unabhängig vom Verbraucherstrom immer konstant. Das bedeutet, daß auch der Spannungsabfall an R_V und die Verbraucherspannung konstant bleiben.

3.4.2.4. Allgemeine Hinweise zur Spannungsstabilisierung mit Referenzdioden

Bei der Spannungsstabilisierung mit Referenzdioden sind im einzelnen folgende Gesichtspunkte zu beachten:

- Die spannungsstabilisierende Wirkung einer Z-Diode ist um so besser, je kleiner ihr dynamischer Widerstand r_z im Zenerbereich ist. Den kleinsten Zenerwiderstand besitzen Z-Dioden einer bestimmten Typenserie für Zenerspannungen zwischen 5 und 8 V.
- Bei voller Ausnutzung des Zenerbereichs läßt sich die Spannungsschwankung am Ausgang von Z-Dioden-Schaltungen auf wenige Zehntel Volt begrenzen, selbst, wenn die Eingangsspannung in weiten Grenzen schwankt.

- In Stabilisierungsschaltungen tritt wegen des unbedingt erforderlichen Vorwiderstands zwangsläufig immer ein Spannungsverlust auf. Die Eingangsspannung muß daher immer größer als die Ausgangs- bzw. Verbraucherspannung sein. Als Richtwert gilt: Die Versorgungsspannung sollte etwa doppelt so hoch wie die Verbraucherspannung sein.
- Die für Z-Dioden mit Metallgehäuse angegebenen Herstellerdaten über die zulässige Verlustleistung beziehen sich häufig auf die an ein Kühlblech (z.B. 100 x 100 x 2 mm³ Al) montierte Diode. Das gilt vorwiegend für Z-Dioden ab etwa 1 Watt Verlustleistung. Diese Tatsache ist beim Aufbau und bei der Bemessung von Stabilisierungsschaltungen unbedingt zu berücksichtigen.
- Da die obere Leistungsgrenze handelsüblicher Z-Dioden bei etwa 10 Watt liegt, lassen sich mit Z-Dioden nur Stabilisierungsschaltungen mit verhältnismäßig geringer Strombelastbarkeit (ca. 0,5 A) verwirklichen. Soll die Spannung für Verbraucher größerer Leistungsaufnahme stabilisiert werden, muß man Stabilisierungsschaltungen mit Transistoren anwenden, wie sie im Abschnitt 4.6 beschrieben werden.

3.4.3. Weitere Anwendungsmöglichkeiten für Referenzdioden

Neben ihrem Hauptanwendungsgebiet in Spannungs-Stabilisierungsschaltungen lassen sich Referenzdioden auch für andere Zwecke einsetzen. Dazu drei Beispiele:

Beispiel 1: Spannungsbegrenzung

Obgleich Referenzdioden vorwiegend für den Betrieb im Sperrbereich vorgesehen sind, werden sie vielfach im Durchlaßbereich betrieben. Der Grund dafür liegt im sehr steilen Anstieg des Durchlaßstroms oberhalb der Diffusionsspannung. In Durchlaßrichtung betriebene Referenzdioden dienen vorwiegend als **Spannungsbegrenzer für kleine Spannungen**, z.B. zur Stabilisierung der Basis-Emitter-Spannung von Gegenakt-Verstärkerstufen mit Transistoren. Referenzdioden als Spannungsbegrenzer in Durchlaßrichtung werden in Industrieschaltungen häufig wie einfache Dioden dargestellt. Daß es sich um Referenzdioden handelt, erkennt man dann lediglich aus der Typenbezeichnung (siehe Abschn. 10).

Beispiel 2: Nullpunktunterdrückung

In der Meßtechnik ist es häufig erforderlich, geringe **Spannungsschwankungen** in einem bestimmten Bereich genau zu messen. Nehmen wir z.B. an, eine Gleichspannung schwankt zwischen etwa 24 und 25 V. Diese Spannungsänderung könnte auf der Skala eines 30-V-Zeigerinstruments nur ungenau abgelesen werden, da sich der Zeigerausschlag nur um 1/30 der Skalenlänge ändern würde.

Durch die Reihenschaltung einer Referenzdiode zum Meßinstrument läßt sich nun der nicht interessierende Spannungsbereich unterdrücken. Wir wählen eine Z-Diode mit einer Zenerspannung von 22 V und schalten in Reihe dazu ein Voltmeter mit dem Meßbereich 3 V. Der Zeiger des Spannungsmessers wird erst ausschlagen, wenn die Meßspannung $U_{\text{Meß}}$ größer ist als die Zenerspannung (22 V); denn bis zu dieser Spannung ist die Z-Diode so hochohmig (Megohm), daß praktisch kein Strom über das Instrument fließt. Auf der Skala unseres 3-V-Spannungsmessers müßte also jetzt die Zeigerruhelage mit 22 V und der Zeigervollauschlag mit 25 V geeicht werden. Da die Skala nicht beim Wert Null beginnt, wird diese Schaltungsmaßnahme als **Nullpunktunterdrückung** bezeichnet.

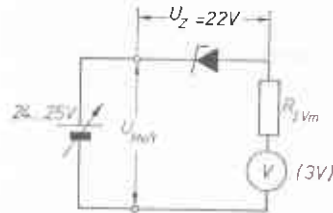


Abb. 3.52 — Nullpunktunterdrückung bei einem Voltmeter

In unserem Beispiel ergibt jetzt eine Spannungsänderung von 24 auf 25 V eine Änderung des Zeigerausschlags um $\frac{1}{3}$ im 3-V-Skalenbereich und kann daher 10mal genauer abgelesen werden als auf der 30-V-Skala.

Beispiel 3: Brummsieb

Eine Referenzdiode kann aufgrund ihrer Begrenzerwirkung auch als **Brummsieb** hinter Netzgleichrichtern eingesetzt werden. Schaltet man hinter einen Brückengleichrichter entsprechend nebenstehender Abbildung eine Referenzdiode, so hat sie die gleiche Wirkung wie ein Siebkondensator. Die Zenerspannung der Referenzdiode soll etwa gleich dem Effektivwert der gleichgerichteten Wechselspannung oder etwas kleiner sein. Da der Spitzenwert einer sinusförmigen Spannung um den Faktor 1,414 größer ist als der Effektivwert, begrenzt die Referenzdiode die gleichgerichteten Spannungsimpulse. Die Sinushalbwellen

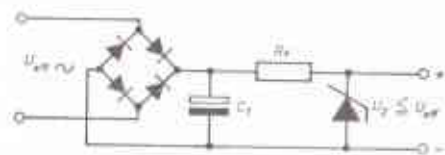


Abb. 3.53 — Referenzdiode als Brummsieb

werden somit stark abgeflacht und die Welligkeit der pulsierenden Gleichspannung herabgesetzt. Im zeitlichen Verlauf der pulsierenden Gleichspannung treten infolge der Begrenzerwirkung erheblich kürzere Lücken mit Spannungseinbruch auf. Diese kurzen Lücken werden durch den Ladekondensator C_1 fast vollständig überbrückt. Das nebenstehende Spannungs-Zeit-Diagramm zeigt nur noch ganz kurzzeitige Spannungsschwankungen.

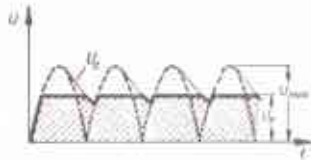


Abb. 3.54 — Wirkung der Referenzdiode als Brummsieb

Die Referenzdiode als Brummsieb ist mit einem Siebkondensator der Kapazität

$$C = \frac{1}{\omega \cdot r_x} \quad \text{vergleichbar,}$$

wobei vorausgesetzt wird, daß $U_z \leq U_{\text{eff}}$ ist.

Beträgt also beispielsweise der Zenerwiderstand der Referenzdiode 2Ω und die Frequenz der Brummspannung 100 Hz (Brückenschaltung), so ergibt sich ein Kapazitätswert von

$$C = \frac{1}{2\pi \cdot 100 \cdot 2} \quad F = \frac{1}{1256} \quad F = 0,000796 \text{ F} \approx 800 \mu\text{F}.$$

Neben dem geringen Raumbedarf gegenüber einem Kondensator gleicher Kapazität bietet die Referenzdiode hier noch den Vorteil, daß sie gleichzeitig spannungsstabilisierend wirkt.

3.5. Varistoren (VDR-Widerstände)

Varistoren sind Halbleiter-Bauelemente, deren Widerstand richtungsunabhängig mit zunehmender Spannung kleiner wird. Die Bezeichnung ist aus dem englischen Begriff *variable resistor* (= veränderbarer Widerstand) entstanden. Daneben findet man sehr häufig die Bezeichnung VDR-Widerstand, abgeleitet aus *voltage dependent resistor* (= spannungsabhängiger Widerstand).

3.5.1. Grundsätzlicher Aufbau und Wirkungsweise

Die Varistoren bestehen aus feinkörnigem Siliziumkarbid (SiC), das unter mehr oder weniger hohem Druck zu überwiegend scheibenförmigen Widerstandskörpern zusammengesintert wird. An den Berührungsflächen der nebeneinanderliegenden Siliziumkarbidkörner tritt eine Gleichrichterwirkung auf (PN-Übergang). Da in einem Varistor aber auf kleinstem Raum sehr viele Körner nebeneinander und deren Berührungsflächen räumlich ganz willkürlich zueinander liegen, ergibt sich keine eindeutige Sperr- oder Durchlaßrichtung. Man kann sich einen Varistor etwa so vorstellen, wie es die Abb. 3.55 zeigt.

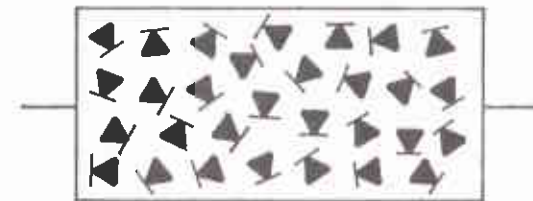


Abb. 3.55 — Schematischer Aufbau eines Varistors
(zusammengesetzt aus vielen Dioden mit willkürlicher Lage)

Legt man an ein solches Bauelement eine Spannung bestimmter Richtung, so kann der Strom nur über die Kette der zufällig in Durchlaßrichtung liegenden Diodenstrecken fließen. Wir wissen aber von

einem PN-Übergang, daß über ihn in Durchlaßrichtung erst ein Strom fließt, wenn die angelegte Spannung größer als die Diffusionsspannung ist. Da ein Varistor eine Vielzahl von PN-Übergängen in Reihe geschaltet ist, addieren sich deren Diffusionsspannungen. So kommt es, daß erst bei verhältnismäßig hohen Spannungen ein merklicher Stromanstieg erfolgt. Wird die angelegte Spannung umgepolt, so liegt wieder eine andere Kette von PN-Übergängen zufällig in der Durchlaßrichtung. In entgegengesetzter Richtung wiederholt sich der gleiche Vorgang. Erst wenn die angelegte Spannung größer als die Summe der Diffusionsspannungen ist, beginnt der Strom zu fließen.

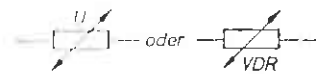


Abb. 3.56

Als Schaltzeichen für Varistoren finden die nebenstehenden Symbole Verwendung, von denen das linke genormt, das rechte jedoch häufiger in Schaltungen zu finden ist.

Das Verhalten eines Varistors wird durch seine Kennlinie anschaulich; sie zeigt einen symmetrischen Verlauf für positive und negative Spannungsrichtung.

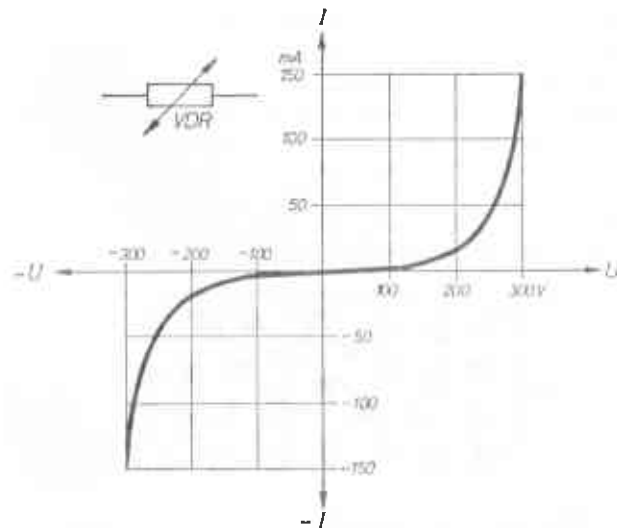


Abb. 3.57 — Kennlinie eines Varistors (VDR-Widerstands)

Es liegt nahe, daß man die Spannung, bei der ein Stromanstieg erfolgt, durch die Zahl der PN-Übergänge beeinflussen kann. Verwendet man zur Herstellung von VDR-Widerständen grobkörniges Sili-

ziumkarbid, so sind auf gleichem Raum weniger PN-Übergänge vorhanden. Die Summe der Diffusionsspannungen wird kleiner. Der Strom steigt daher schon bei einer geringeren angelegten Spannung. Ein Varistor besteht aus einer Vielzahl willkürlich zueinanderliegender PN-Übergänge. Sowohl bei positiver als auch bei negativer Richtung der angelegten Spannung beginnt ein merklicher Stromanstieg erst bei einer bestimmten Spannungshöhe. Der Vorteil der VDR-Widerstände liegt darin, daß sie gegenüber Begrenzer- und Referenzdioden richtungsunabhängig sind. Varistoren werden für Spannungen zwischen etwa 10 und 1000 Volt hergestellt.

3.5.2. Anwendung von Varistoren

Wie die Erläuterung der grundsätzlichen Wirkungsweise schon vermuten läßt, werden Varistoren vorwiegend zur Begrenzung von Wechselspannungen bzw. Pulsspannungen oder Spannungsschößen eingesetzt.

Beispiel 1: Funkenlöschung

Beim Unterbrechen eines induktiv belasteten Stromkreises entstehen sehr hohe Selbstinduktionsspannungen. Diese Induktionsspannungen betragen ein Vielfaches der Betriebsspannung und können zum Durchschlag (Zerstörung von Wicklungen oder Bauelementen), zu starker Funkenbildung an Schalterkontakten (Abbrand) und Geräuschstörungen (Knacken) führen. Die auftretenden Spannungsspitzen lassen sich durch einen Varistor einfach unterdrücken. Dazu ist folgender Versuch anschaulich, der auch zur Demonstration der Selbstinduktionsspannung an einer Spule gebräuchlich ist:

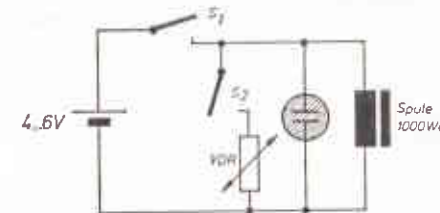


Abb. 3.58 — Funkenlöschung mit einem Varistor

Eine Spule (ca. 1000 Windungen) mit Eisenkern wird über einen Schalter s_1 an eine Gleichspannung von etwa 4 bis 6 Volt gelegt. Parallel zur Spule ist eine Glimmlampe mit einer Zündspannung von 60 bis 70 Volt geschaltet. Solange der Stromkreis geschlossen ist, leuchtet die Glimmlampe nicht; denn die anliegende Spannung liegt weit unter der erforderlichen Zündspannung. Wird der Schalter s_1 jedoch geöffnet, so blitzt die Glimmlampe kurzzeitig auf. Das beweist, daß die Selbstinduktionsspannung mindestens 60 bis 70 Volt betragen muß. (Wer die Schalterkontakte zufällig berührt, bekommt einen heftigen Schlag.) Jetzt schalten wir mit Hilfe des Schalters s_2 parallel zur Spule und zur Glimmlampe einen VDR-Widerstand, der bei etwa 20 ... 30 V niederohmig wird. Beim Öffnen des Schalters s_1 leuchtet die Glimmlampe nicht mehr auf. Für die hohe Selbstinduktionsspannung stellt der VDR-Widerstand einen niederohmigen Nebenschluß dar. Dieses Verfahren, hohe Spannungsspitzen beim Öffnen eines Schalters zu unterdrücken, ist als **Funkenlöschung** bekannt.

Beispiel 2: Überspannungsschutz

Ähnlich — wie im Beispiel 1 — liegen die Verhältnisse in induktiv belasteten Stromkreisen, wenn mit Pulsspannungen, z.B. Nadelimpulsen, oder auch Kippströmwindungen (Oszilloskop bzw. Fernsehempfänger) gearbeitet wird. Da in diesen Fällen eine sehr schnelle Änderung des magnetischen Flusses in einer Induktivität auftritt, entstehen wie beim Schalten sehr hohe Selbstinduktionsspannungen. Um die Spulen selbst oder auch andere Bauelemente des Stromkreises (Röhren, Transistoren, Kondensatoren) vor einer Zerstörung zu schützen, werden Varistoren eingeschaltet. Nehmen Sie doch einmal die Schaltung eines Fernsehgeräts zur Hand! Sowohl in der Vertikalablenk- als auch in der Horizontalablenkstufe findet man VDR-Widerstände. In Schaltungen, bei denen mit Pulsspannungen gearbeitet wird, dienen VDR-Widerstände häufig auch zum Abflachen der Impulsform.

Neben den aufgezeigten Anwendungsmöglichkeiten für Varistoren bietet sich ihre Verwendung zur Stabilisierung von Spannungen (ähnlich wie Z-Dioden), zur Impulsformung und zum allgemeinen Überspannungsschutz an.

3.6. Fotodioden

Fotodioden gehören genaugenommen in die Gruppe der durch äußere Einflüsse (Wärme, Magnetfeld, Licht) veränderbaren Widerstände. Da sie jedoch als wesentliches Merkmal einen PN-Übergang enthalten, unterscheiden sie sich in ihrer Wirkungsweise von Fotowiderständen dadurch, daß sie stromrichtungsabhängig sind.

3.6.1. Grundsätzlicher Aufbau und Wirkungsweise

Eine Fotodiode enthält einen aus entsprechend dotiertem Germanium oder Silizium gebildeten PN-Übergang. Die Halbleiterschichten sind bei ihr jedoch so dünn gehalten, daß sie vom einfallenden Licht durchsetzt werden können. Das Licht kann also bis zum eigentlichen PN-Übergang vordringen. Um eine Beeinflussbarkeit durch Fremdlicht zu vermeiden, sind sie in lichtundurchlässige Gehäuse aus Kunststoff oder Metall eingebaut, die nur eine kleine Öffnung enthalten. In die Öffnung wird meistens eine Glaslinse zur Lichtbündelung eingeschmolzen. Die lichtempfindliche Fläche beträgt nur etwa 1 mm^2 .



Abb. 3.59 — Germanium-Fotodioden

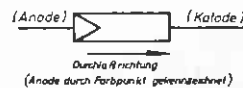


Abb. 3.60 — Schaltzeichen für Fotodioden

Eine Fotodiode wird grundsätzlich in Sperrichtung betrieben. Wird eine Spannung in Sperrichtung angelegt, so fließt ein sehr geringer

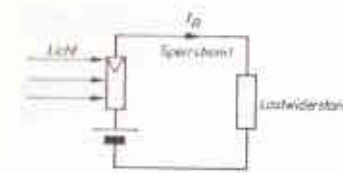


Abb. 3.61 — Schaltung einer Fotodiode

auf Eigenleitfähigkeit beruhender Sperrstrom (Dunkelstrom). Sobald aber Licht auf die Sperrschicht fällt, werden in ihr Ladungsträger (Elektronen) freigesetzt, wodurch sich der Strom in Abhängigkeit von der Lichtstärke um ein Vielfaches gegenüber dem Sperrstrom im unbeleuchteten Zustand erhöht. Die Stromstärke ist dabei nahezu unabhängig von der Höhe der angelegten Spannung; sie steigt jedoch fast linear mit der Beleuchtungsstärke. Daher sind Fotodioden für meßtechnische Zwecke (Lichtmessung) sehr gut geeignet. Nachteilig gegenüber Fotowiderständen ist allerdings die geringe Strombelastbarkeit (Größenordnung $100\ \mu\text{A}$) und die geringe zulässige Verlustleistung (20 bis 100 mW). Fotodioden eignen sich nicht zur direkten Steuerung von Relais, sondern sind nur in Verbindung mit empfindlichen Meßgeräten oder Verstärkern anwendbar. Dagegen besitzen sie aber noch einen wesentlichen Vorteil gegenüber Fotowiderständen: die hohe Grenzfrequenz. Während sich Fotowiderstände nur für Wechselvorgänge bis zu 3 Hz eignen, reagieren Fotodioden noch auf Frequenzen bis zu einigen 100 000 Hz.

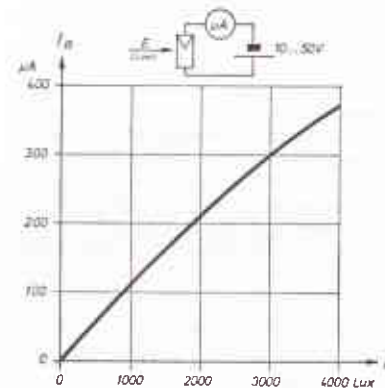


Abb. 3.62 — Abhängigkeit des Sperrstroms einer Fotodiode von der Beleuchtungsstärke

Ihre größte spektrale Empfindlichkeit liegt im Bereich des Infrarotlichts. Für sichtbares Licht sind Si-Fotodioden empfindlicher als Ge-Fotodioden. Da Fotodioden sehr klein gebaut werden können (Durchmesser ca. 1 mm !), sind sie für den Einsatz in elektronischen Schaltungen hervorragend geeignet.

Zur Abtastung von Lochstreifen werden von der Industrie auch komplette Fotodiodenstreifen — bestehend aus 5 oder mehr Fotodioden — hergestellt.

Fotodioden werden in Sperrichtung betrieben. Bei Belichtung steigt der Sperrstrom nahezu linear mit der Beleuchtungsstärke. Die zulässi-

ge Verlustleistung ist gering. Sie sind für Wechselvorgänge bis ca. 100 kHz geeignet.

Die Fotodiode läßt sich übrigens auch als Fotoelement verwenden. Liegt keine Sperrspannung an, so wird in der Fotodiode mit zunehmender Beleuchtungsstärke eine Spannung erzeugt. Diese Spannung hat bei 100 Lux Beleuchtungsstärke eine Höhe von etwa 100 bis 400 Millivolt und ist groß genug, um beispielsweise einen Germaniumtransistor durchzusteuern.

3.6.2. Anwendung von Fotodioden

Der breiteste Anwendungsbereich für Fotodioden besteht in der Abtastung schneller optischer Vorgänge. Dazu gehört u.a. auch die Abtastung des Lichttonstreifens bei der Tonfilmwiedergabe (früher verwendete man dafür Vakuum-Fotozellen).

Zur Anwendung von Fotozellen in der Datenverarbeitung ein Beispiel: Informationen, die einem Datenverarbeitungsgerät eingegeben werden sollen, sind vielfach in Form von gestanzten Löchern auf Karten oder Papierstreifen gespeichert. Für die Schnelligkeit der Datenverarbeitung ist u.a. maßgebend, mit welcher Geschwindigkeit die Lochkarten oder -streifen „gelesen“ werden können. Eine Fotodiode ermöglicht wegen ihrer hohen Grenzfrequenz die Abtastung von bis zu 100 000 Löchern pro Sekunde. Man läßt dazu einfach die Lochkarten oder -streifen eine Lichtschranke mit einer Fotodiode als Lichtempfänger passieren. Sind — wie auf einer Lochkarte — mehrere Lochreihen vorhanden, so können wegen ihrer geringen Baugröße entsprechend der Reihenzahl mehrere Fotodioden zu einer sogenannten Zeile zusammengefaßt werden. Beim Durchlaufen der Lochkarte fällt durch die Löcher kurzzeitig Licht auf die Fotodioden und löst einen Stromimpuls aus.

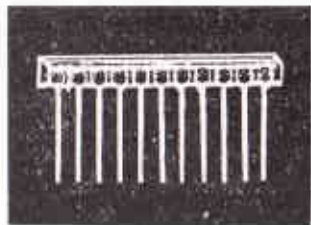


Abb. 3.63 — Zelle mit 10 Fotodioden

Nach einem ähnlichen Verfahren werden die in Form von Punkten codierten Postleitzahlen auf Briefen und Karten bei der elektronischen Briefverteilung abgetastet. Die bei normalem Licht unsichtbaren Punkte leuchten bei Bestrahlung mit Ultraviolettlicht auf.

1.7. Lumineszenzdiolen (LED)

1.7.1. Grundsätzlicher Aufbau und Wirkungsweise

Die Steuerung elektronischer Schaltungen durch Licht gewinnt immer mehr an Bedeutung. Glühlampen weisen jedoch als Lichtquellen in der modernen Elektronik eine Reihe von Nachteilen auf. Zu den Nachteilen der Glühlampen gehören eine viel zu große Trägheit in bezug auf die geforderten kurzen Schaltzeiten (ns), die geringe Lebensdauer von größenordnungsmäßig etwa 1000 Brennstunden, die für fotoelektronische Empfänger ungünstige Wellenlänge des Glühlampenspektrums und die nichtlineare Helligkeitssteuerung.

In den letzten Jahren wurden von der Halbleiterindustrie sogenannte Lumineszenzdiolen entwickelt, die die für die Elektronik geforderten Bedingungen als Lichtsender weitaus besser erfüllen als Glühlampen. Lumineszenzdiolen — auch LED (= lichtemittierende Diolen) genannt — werden aus den Grundmaterialien Galliumphosphid (GaP) oder Galliumarsenphosphid (GaAsP) oder Galliumarsenid (GaAs) hergestellt. Durch Eindiffundieren einer P-Zone entsteht ein PN-Übergang. Wird an den PN-Übergang eine Spannung von ca. 1,3 V bis 2 V (je nach Typ) in Durchlaßrichtung gelegt, so entsteht infolge Rekombination von Elektronen mit Löchern eine Lichtstrahlung. Die Stromstärke liegt dabei zwischen etwa 10 mA und 100 mA.

Da Elektronen nur durch Energiezufuhr von einer Schale entfernt werden können, geben sie diese zusätzlich aufgenommene Energie wieder ab, wenn sie auf eine gleichartige Schale zurückkehren. Die dabei freiwerdende Energie kann je nach Aufbau der Materie z.B. Wärme oder Licht sein.

Die Farbe des von einer Lumineszenzdiode ausgestrahlten Lichts hängt vom verwendeten Ausgangsmaterial ab. In Abb. 3.64 sind für drei verschiedene typische Lumineszenzdiolen die Spektralkurven dargestellt. Spektralkurven geben an, wie stark eine Lichtquelle bei einer bestimmten Lichtwellenlänge (Farbe) leuchtet.

Das Diagramm läßt erkennen, daß die Lichtfarbe der GaP-Diolen grün-gelb, der GaAsP-Diolen rot-orange und die der GaAs-Diolen infrarot ist. Durch Legierung in einem anderen Mischungsverhältnis lassen sich die Lichtfarben noch in gewissen Grenzen verändern.

Um den zweckmäßigen Einsatz der verschiedenen Lumineszenzdiolen erkennen zu können, sind im Diagramm ebenfalls die Empfindlichkeitskurven für das menschliche Auge und für Silizium-Fotodioden angetragen. Da das menschliche Auge für gelbes Licht die höchste Emp-

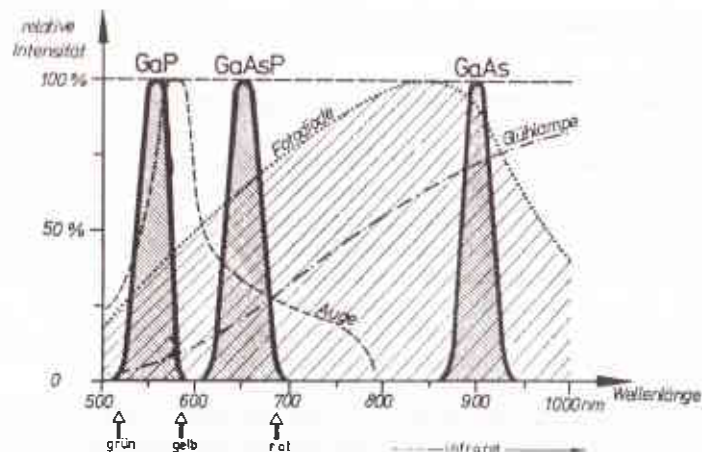


Abb. 3.64 — Spektralkurven einiger Lumineszenzdioden

findlichkeit besitzt, sind für optische Signale Galliumphosphid-Dioden (grün-gelb) am besten geeignet. Dagegen sind z.B. für Lichtschranken mit Fotodioden Galliumarsenid-Dioden (infrarot) vorzuziehen. Interessant ist zum Vergleich die Spektralkurve des Glühlampenlichts. Sie zeigt, daß der überwiegende Anteil des Glühlampenlichts außerhalb des Empfindlichkeitsbereichs des menschlichen Auges oder der Fotodioden liegt. Die Spektralkurven der LED sind so schmal, daß man von praktisch einfarbigem Licht sprechen kann. So lassen sich ganz nach Bedarf durch bestimmte Lumineszenzdioden ganz bestimmte Lichtfarben erzielen, was bei Glühlampen nur mit besonderen Farbfiltern und schlechtem Wirkungsgrad erreicht werden kann.

Die Lebensdauer der Lumineszenzdioden beträgt etwa eine Million (10^6) Stunden gegenüber Tausend (10^3) Stunden einer Glühlampe. **Lumineszenzdioden sind so trägheitsarm, daß sie sich noch für Wechselvorgänge bis in die Größenordnung Megahertz anwenden lassen.** Ihre Anstiegszeit (= Zeit zwischen Einschalten und voller Helligkeit) beträgt nur Nanosekunden. Der Wirkungsgrad ist mit rd. 4 % zwar gering, jedoch im Vergleich mit der Glühlampe noch recht gut. Die Verlustleistung liegt zwischen 50 mW und einigen Watt. LED sind besonders empfindlich gegen zu hohe Sperrspannungen (ca. 3 V) und müssen daher in vielen Anwendungsbereichen durch besondere Schaltungsmaßnahmen vor zu großer Sperrspannung geschützt werden, was sich vielfach durch Begrenzung mit einer Si-Diode erzielen läßt. Zu den besonderen Eigenschaften der Lumineszenzdioden muß auch ihre **lineare Helligkeitssteuerung** gerechnet werden; d.h., daß ihre Helligkeit praktisch linear mit der Durchlaßstromstärke steigt (Abb. 3.65). Damit lassen sie sich für feine Regel- und Steuervorgänge anwenden.

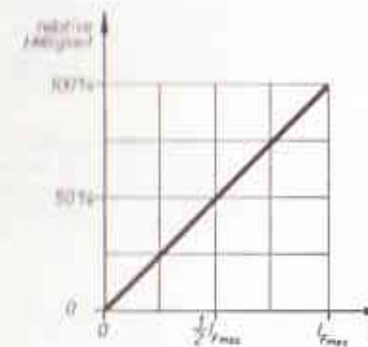


Abb. 3.65 — Abhängigkeit der Helligkeit von der Durchlaßstromstärke

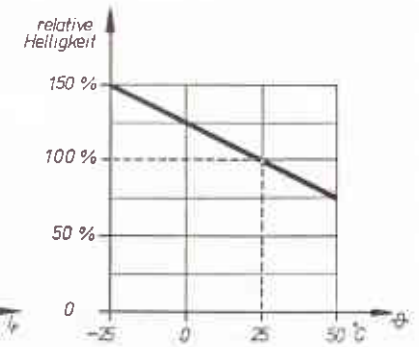


Abb. 3.66 — Abhängigkeit der Helligkeit von der Temperatur

Wie alle Halbleiter-Bauelemente zeigen leider auch Lumineszenzdioden eine starke Temperaturabhängigkeit. Das Diagramm in Abb. 3.66 läßt erkennen, daß die Helligkeit bei einer Temperaturabnahme um 25 °C auf rd. 125 % steigt, dagegen bei einer Temperaturzunahme um 25 °C auf rd. 75 % des Werts bei 25 °C fällt.

Die Abb. 3.67 gibt zwei typische Bauformen von Lumineszenzdioden wieder. Bei der einfachen Leuchtdiode (linke Abbildung) wird der austretende Lichtstrahl vielfach durch eine eingeschmolzene Linse gebündelt. Rechts ist die Kombination mehrerer Lumineszenzdioden zu einem **Siebensegment-Ziffernanzeige-Baustein** dargestellt.



Abb. 3.67 — Bauformen von Lumineszenzdioden

Jedes der sieben Segmente enthält 5 Leuchtdioden und kann einzeln von außen angesteuert werden. Durch entsprechende Kombination der sieben Segmente lassen sich die Ziffern von 0 bis 9 darstellen, wie es unter der rechten Abbildung erkennbar ist.

Als Ersatz für Glühlampen werden — z.Z. noch vornehmlich in USA — Lumineszenzdioden mit eingebautem Vorwiderstand als sogenannte **Patronenlampen** hergestellt. Die Betriebsspannungen betragen 4 ... 6 V, 7 ... 14 V oder 15 ... 30 V bei ca. 20 mA Stromaufnahme. Aus Preisgründen ist ihr Einsatz allerdings z.B. noch auf Sonderfälle beschränkt.

3.7.2 Anwendung von Lumineszenzdioden

Aus dem äußerst breiten Anwendungsbereich hier einige Beispiele:

Beispiel 1: Kontrolllampen

Da sie äußerst klein hergestellt werden können, lassen sich Lumineszenzdiolen an praktisch jeder gewünschten Stelle unmittelbar in eine gedruckte Schaltung einfügen. Damit ermöglichen sie die Überwachung der Funktionsfähigkeit bestimmter Schaltungsteile. Ebenso kann eine Vielzahl solcher „Kontrolllampen“ auf kleinstem Raum an der Frontplatte eines Geräts untergebracht werden.

Beispiel 2: Anzeigelampen

In batteriebetriebenen Geräten mußten zur Abstimm- oder Aussteuerungsanzeige entweder Leuchtöhren (mag. Band) oder Drehspulinstrumente verwendet werden. Beide bedingen aber einen ziemlich hohen Raumbedarf und Schaltungsaufwand. Lumineszenzdiolen können mit niedrigen Spannungen (Transistorgeräte) betrieben werden und sind so klein, daß ihr Raumbedarf nicht ins Gewicht fällt. Allerdings läßt sich mit Hilfe von Lumineszenzdiolen nur eine verhältnismäßig grobe Helligkeitsanzeige bewirken, weil das menschliche Auge gegenüber Helligkeitsunterschieden nicht sehr empfindlich ist.

Beispiel 3: Anzeige von Ziffern und Zeichen

Siebensegment-Ziffernanzeige-Bausteine mit Leuchtdiolen ersetzen die Ziffernanzeigeröhren, die nach dem Glüh- oder Glühlampenprinzip arbeiten. Die Anzeige durch Lumineszenzdiolen ist schneller und kontrastreicher. Es gibt Bausteine mit Ziffern- oder Zeichen.

Beispiel 4: Galvanische Trennung von Stromkreisen

In der Steuer- und Regeltechnik ist es oft wünschenswert, daß Steuerkreis und gesteuerter Stromkreis galvanisch voneinander getrennt sind. Wegen ihrer Trägheitsarmut lassen sich Lumineszenzdiolen als steuernde Lichtquellen zur Übertragung von Wechselvorgängen verwenden. Fotodiolen dienen dann als Lichtempfänger. Lichtsender und -empfänger können galvanisch völlig getrennt sein, da das Licht als Koppel-element wirkt. Für solche Zwecke werden von der Industrie schon Bausteine angeboten, die eine Leucht- und eine Fotodiode enthalten.

Zur Demonstration der Lichtübertragung von Wechselvorgängen eignet sich eine verhältnismäßig einfache Schaltung nach Abbildung 3.68, in der eine infrarotstrahlende Lumineszenzdiode mit einer Verlustleistung von $P_V \approx 200 \text{ mW}$ verwendet wird. Als Wechselstromgenerator dient ein NF-Verstärker mit angeschlossener Tonquelle. Ein Übertrager U sorgt für die galvanische Trennung zwischen Verstärker- und Lumineszenzdiolen-Stromkreis. Der Vorwiderstand R_V ermöglicht die richtige Arbeitspunkteinstellung. Die Lumineszenzdiode wird durch eine Siliziumdiode vor zu hoher Sperrspannung geschützt; denn für die Sperrspannung der LED liegt die Si-Diode in

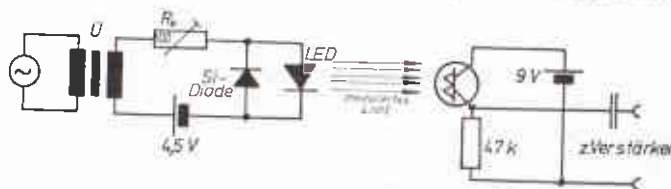


Abb. 3.68 — Tonübertragung mit moduliertem Licht

Durchlaßrichtung und begrenzt sie auf ca. 0,7 V. Damit der Arbeitspunkt in der Mitte des geradlinigen Kennlinienteils liegt, muß die LED durch eine Gleichspannung vorgespannt werden. Sonst ist mit starken Verzerrungen bei der Übertragung zu rechnen. Der Arbeitspunkt wird mit R_V so eingeregelt, daß im Ruhezustand höchstens $\frac{1}{2} I_{Fmax}$ fließen kann.

Als Empfänger dient ein rauscharmer Fototransistor in Kollektorschaltung (Impedanzwandler). Er dient nur als Wandler und muß an einen NF-Verstärker angeschlossen werden. Das von der Lumineszenzdiode ausgestrahlte Licht wird mit Hilfe des Verstärkers moduliert; d.h., ihre Helligkeit schwankt im Rhythmus der Tonfrequenz. Im „Empfänger“ mit nachgeschaltetem Verstärker ist das übertragene Signal wieder als Ton hörbar. Die Wirksamkeit der Lichtübertragung läßt sich einfach dadurch beweisen, daß man eine Hand zwischen „Sender“ und „Empfänger“ hält; der Empfänger ist dann stumm. Die Reichweite dieser Versuchsanordnung beträgt allerdings nur wenige Meter. Für größere Reichweiten sind besonders leistungsfähige Lumineszenzdiolen und eine gute Bündelung des Lichts durch Hohlspiegel bei Sender und Empfänger anzuwenden.

Nach diesem Prinzip arbeiten Lichtsprengeräte und sinngemäß auch die Nachrichtenübertragung mit Laserstrahlen.

Beispiel 5: Lichtquelle in Lichtschranken

Da das Lichtspektrum einer Lumineszenzdiode äußerst schmal ist, kann man mit ihr fast einwelliges Licht erzeugen. Eine solche Lichtquelle bietet für Lichtschranken neben der Diebstahlsicherung die Möglichkeit zur optischen Sortierung — beispielsweise für die Parbertlichkeit von Erzeugnissen.

3.8. Kapazitätsdiolen (Varaktoren)

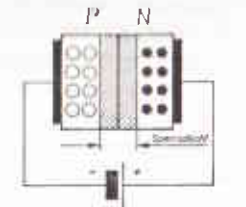


Abb. 3.69 — In Sperrrichtung betriebener PN-Übergang (schematisch)

Ein in Sperrrichtung betriebener PN-Übergang wirkt wie eine Kapazität. Betrachten wir die schematische Darstellung eines gesperrten PN-Übergangs, so zeigt sich folgendes: Zwischen zwei leitenden Schichten (P- und N-Leiter) ist eine Sperrschicht vorhanden, die praktisch frei von Ladungsträgern ist. Diese Sperrschicht bildet das Dielektrikum, die P- und die N-Zone die beiden Beläge eines Kondensators.

Die Kapazität eines Kondensators hängt nach der Formel $C = \frac{A \cdot \epsilon}{d}$ von der Plattenfläche A, der Dielektrizitätskonstante ϵ und dem Plattenabstand d ab. Bei einer Kapazitätsdiode liegt die Plattenfläche durch die Baugröße und die Dielektrizitätskonstante durch das verwendete Material (Germanium oder Silizium) fest. Der Plattenabstand

— hier die Dicke der Sperrschicht — läßt sich aber leicht durch die Höhe der Sperrspannung beeinflussen; es bedeuten:

**Geringe Sperrspannung = dünne Sperrschicht, also große Kapazität,
hohe Sperrspannung = dicke Sperrschicht, also kleine Kapazität.**

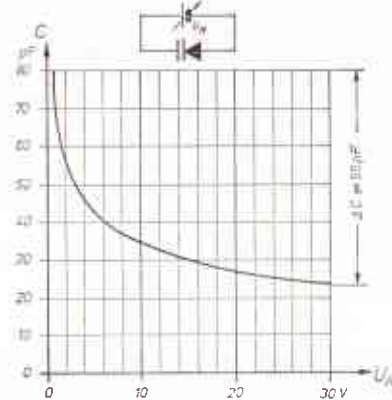


Abb. 3.70 — Abhängigkeit der Kapazität von der Höhe der Sperrspannung

lauf. Der Kennlinienverlauf einer Kapazitätsdiode wird auch noch dadurch beeinflusst, daß die elektrische Feldstärke mit zunehmendem Plattenabstand geringer wird.

Da sich die Halbleiterstoffe Germanium und Silizium wegen ihrer großen Härte und Sprödigkeit nicht zu größeren Bauelementen verarbeiten lassen, ist die wirksame Plattenfläche einer Kapazitätsdiode verhältnismäßig gering. Größere Kapazitätswerte lassen sich daher nur durch kräftiges Dotieren und geringe Schichtdicke (geringen Abstand) der Sperrzone erreichen. Die untere Kapazitätsgrenze ist durch die Höhe der zulässigen Sperrspannung gegeben. Für Kapazitäts-

dioden gibt es an sich kein genormtes Schaltzeichen, sondern nur das übliche Diodensymbol. Einige Firmen sind jedoch dazu übergegangen, die Kapazitätsdioden in Schaltungen besonders zu kennzeichnen; am häufigsten werden nebenstehende Schaltzeichen verwendet. Kapazitätsdioden können — gleicher Kapazitäts-Änderungsbereich vorausgesetzt — Kondensatoren mit mechanisch veränderbarer Kapazität (z.B. Drehkondensatoren) ersetzen. Die Baugröße ist erheblich geringer als

die der Drehkondensatoren. Da zur Beeinflussung der Kapazität die Höhe der Sperrspannung verändert wird, ist sogar eine einfache Fernregelung über beliebig lange Leitungen möglich. Resonanzkreise in

Die Abhängigkeit von der Sperrspannung ist — wie aus dem Diagramm Abb. 3.70 ersichtlich — jedoch nicht linear.

Das läßt sich leicht erklären: Mit zunehmender Sperrspannung wird die an Ladungsträgern verarmte Sperrzone am PN-Übergang dicker. In der Formel zur Kapazitätsberechnung steht aber der Plattenabstand d im Nenner. Da Plattenfläche A und Dielektrizitätskonstante ϵ für eine bestimmte Diode konstant sind, kann man auch schreiben

$$C \sim \frac{1}{d}$$

Daraus folgt, daß der doppelte

Plattenabstand die Kapazität auf die Hälfte, der vierfache Abstand die Kapazität auf ein Viertel des ursprünglichen Werts herabsetzt. Zeichnet man ein entsprechendes Diagramm, so zeigt es den für die Kennlinie einer Kapazitätsdiode typischen Verlauf.

Funksende- und -empfangsgeräten werden in zunehmendem Maße durch Kapazitätsdioden abgestimmt.

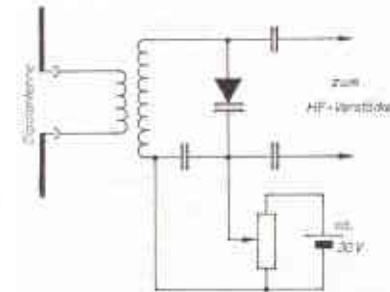


Abb. 3.72 — Abstimmung eines Hochfrequenz-Resonanzkreises mit einer Kapazitätsdiode

In Abb. 3.72 ist eine einfache Schaltung zur Abstimmung eines Parallelresonanzkreises mit einer Kapazitätsdiode dargestellt. Mit dem Potentiometer kann die Sperrspannung an der Kapazitätsdiode und damit ihre Kapazität verändert werden. Der Resonanzkreis ist so leicht auf die gewünschte Frequenz abzustimmen; die übrigen Kapazitäten dienen nur zur gleichstrommäßigen Trennung. Diese Schaltung läßt sich noch erweitern, indem man mehrere Potentiometer verwendet, von denen jedes auf einen gewünschten Wert (Sender) eingestellt und mit einem Schalter an die

Kapazitätsdiode gelegt werden kann (z.B. Sendertasten eines Fernsehempfängers mit Diodentuner). Damit sich die Abstimmung nicht verändern kann, muß die Versorgungsspannung unbedingt stabilisiert werden. Da jedoch die Kapazitätsdioden in Sperrrichtung betrieben werden, ist der Strombedarf sehr gering. Für die Spannungsstabilisierung werden daher nur Referenzdioden mit geringer Verlustleistung benötigt.

4. Der Transistor

4.1. Grundsätzlicher Aufbau und Wirkungsweise

Um die Wirkungsweise von Halbleiter-Bauelementen zu verstehen, ist es zweckmäßig, mit einfachen Vorstellungen zu arbeiten. Wir wollen daher zum besseren Verständnis für die Betrachtung der Vorgänge in einem Transistor folgende einfache Anschauung zugrunde legen:

In einem N-Leiter sind freie Elektronen die Ladungsträger. Elektronen sind elektrisch negativ geladen. Ein N-Leiter ist also ein Leitermaterial, das frei bewegliche negative Ladungen enthält. Als bildliches Symbol verwenden wir dafür $\bullet$.

Durch die Dotierung eines P-Leiters sind in seinem Gefüge „Löcher“ entstanden. Diese Löcher sind Stellen, an denen zur Kristallbildung Elektronen fehlen. Ein Loch wird daher auch als Defektelektron oder allgemein als Störstelle bezeichnet. Ein Loch hat das Bestreben, ein Elektron, also eine negative Ladung einzufangen. Für unsere Begriffe müssen aber Stoffe, die Elektronen selbständig aufnehmen, selbst elektrisch positiv geladen sein. Somit können wir einen P-Leiter als ein Leitermaterial ansehen, das freie positive Ladungen enthält. Dafür wollen wir das Symbol $\oplus$ verwenden.

Bei unseren künftigen Betrachtungen müssen wir ferner davon ausgehen, daß zwischen der Bewegungsrichtung der freien Elektronen in einem Leiter und der technischen Stromrichtung zu unterscheiden ist. Die technische Stromrichtung ist entgegengesetzt zur Bewegungsrichtung der Elektronen.

Ein Transistor besteht nun im Gegensatz zu einer Diode aus drei aufeinanderfolgenden Halbleiterschichten. Die Reihenfolge kann sein:

N-Leiter P-Leiter N-Leiter
(man spricht dann von einem NPN-Transistor)

oder auch

P-Leiter N-Leiter P-Leiter
(diese Art wird als PNP-Transistor bezeichnet).

Aus technischen Herstellungsgründen werden Transistoren auf Germaniumbasis überwiegend als PNP-Typen, auf Siliziumbasis dagegen überwiegend als NPN-Typen gefertigt. Die Wirkungsweise ist bei beiden Arten grundsätzlich die gleiche. Da wir bei unseren künftigen Meßversuchen und Schaltungen Siliziumtransistoren verwenden wollen, werden wir die Wirkungsweise des Transistors anhand eines NPN-Typs erläutern. Ein NPN-Transistor besteht — wie bereits erwähnt — aus drei aufeinanderfolgenden Halbleiterschichten (Abb. 4.1).

Die mittlere Schicht muß dabei sehr dünn (praktisch etwa 10...20 μm dick) und schwach dotiert sein. Nach Abb. 4.1 wird die linke N-Schicht

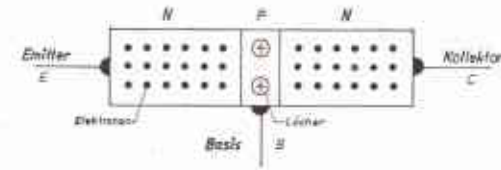


Abb. 4.1 — Schematischer Aufbau eines NPN-Transistors

als **Emitter** bezeichnet (abgeleitet aus dem Lateinischen: emittire = aussenden). Die mittlere Schicht wird **Basis** genannt. Die Bezeichnung ergibt sich aus dem Herstellungsverfahren für Transistoren; Grundmaterial — also Basismaterial — ist dabei fast immer der mittlere Halbleiterstoff. **Kollektor** ist die Bezeichnung für die rechte Schicht (abgeleitet aus dem Lateinischen: colligere = sammeln). Allgemein werden anstelle dieser Bezeichnungen große Buchstaben gesetzt:

Für Emitter **E**,
für Basis **B** und
für Kollektor **C**.

Jede der drei Transistorschichten wird mit einem Anschluß versehen. Die Anschlüsse erhalten die gleiche Benennung, wie die zugehörigen Transistorschichten. Betrachten wir den schematischen Aufbau eines NPN-Transistors, so ist festzustellen, daß in ihm zwei PN-Übergänge vorhanden sind:

- ein PN-Übergang Basis-Emitter und
- ein weiterer PN-Übergang Basis-Kollektor.

Zum Betrieb eines Transistors werden von außen zwei Spannungen angelegt, und zwar so gepolt, wie es Abb. 4.2 zeigt.

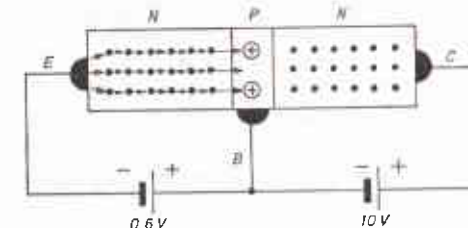


Abb. 4.2 — Grundätzliche Schaltung eines NPN-Transistors

Die Spannung an B—E wollen wir zunächst einmal gerade so groß wählen wie die Diffusions- oder Schwellspannung des PN-Übergangs

Für einen Silizium-PN-Übergang beträgt diese Spannung etwa 0,6 Volt; zwischen Basis und Kollektor wird eine Spannung von etwa 10 Volt angelegt. Wenden wir unsere Kenntnisse über den PN-Übergang an: Der PN-Übergang Basis-Emitter wird in Durchlaßrichtung betrieben. Infolge der angelegten Spannung ist dagegen der PN-Übergang Basis-Kollektor gesperrt.

Was geht nun aufgrund dieser Schaltung in einem Transistor vor? Am PN-Übergang Basis-Emitter liegt eine Spannungsquelle, die mit ihrem Minuspol (Elektronenüberschuß) die negativen Ladungen (Elektronen) des N-Leiters abstößt. Die freien Elektronen der Emitterschicht werden daher in die Basisschicht abgedrängt. Der Emitter sendet also gewissermaßen Elektronen aus (emittieren = aussenden). Da aber die Basiszone sehr dünn und nur schwach dotiert ist, also wenige Löcher enthält, wird auch nur eine geringe Anzahl der in die Basis eingedrungenen Elektronen von den Löchern eingefangen (in der Halbleitertechnik spricht man dabei von „rekombinieren“). Der größte Teil der Elektronen bleibt in der Basis frei beweglich.

Damit ist aber genaugenommen kein PN-Übergang mehr vorhanden; denn auch die Basisschicht enthält ja jetzt ebenso wie die Emitter- und die Kollektorschicht nur freie Elektronen als Ladungsträger. Wir können den Transistor in diesem Zustand mit einem einfachen ohmschen Widerstand vergleichen, in dem ja auch nur freie Elektronen als Ladungsträger vorhanden sind. Für diesen Zustand können wir folgendes vereinfachte Ersatzschaltbild zeichnen, wobei wir die Widerstände E—B und B—C gleich groß annehmen wollen:

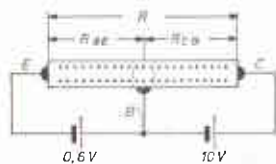


Abb. 4.3

Die in die Basis vorgedrungenen Elektronen können nun in zwei Richtungen abwandern, nämlich sowohl zum Pluspol der Basis-Emitter-Spannungsquelle als auch zum Pluspol der Basis-Kollektor-Spannungsquelle (entgegengesetzte Ladungen ziehen sich an). Da aber das Pluspotential am Kollektor größer ist als an der Basis, fließt der größte Anteil der Elektronen zum Kollektor hin ab. Nur ein geringer Anteil wird vom kleineren Pluspotential an der Basis angezogen.

Zusammenfassend ist hierzu festzustellen, daß infolge der angelegten Spannungen ein großer Elektronenstrom über den Kollektoranschluß fließt. Demgegenüber fließt jedoch nur ein sehr kleiner Elektronenstrom über den Basisanschluß des Transistors.

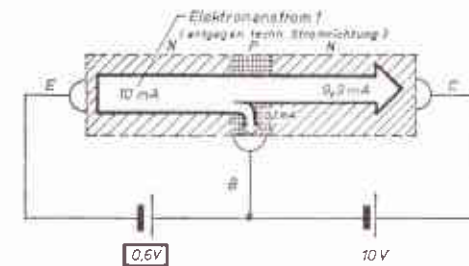


Abb. 4.4 — Stromverteilung in einem Transistor

Jetzt wollen wir die Spannung zwischen Emitter und Basis etwas verringern, nämlich auf etwa 0,5 Volt. Was ergibt sich daraus? Bei herabgesetzter Spannung fließen entsprechend weniger Elektronen in die Basisschicht. Wenn aber die Basisschicht wenige freie Elektronen enthält, können auch nur wenige Elektronen zum Kollektor und zur Basis abfließen. Wird die Spannung so weit verringert, daß sie viel kleiner als die Diffusionsspannung des PN-Übergangs B—E ist, so fließen schließlich so gut wie keine Elektronen mehr; der PN-Übergang ist jetzt sehr hochohmig.

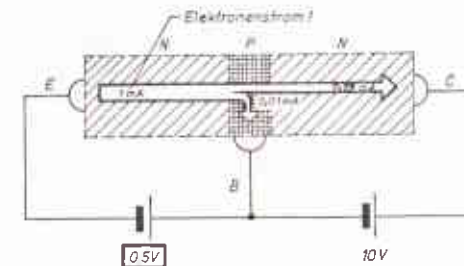


Abb. 4.5 — Stromverteilung bei verringerter Basis-Emitter-Spannung

Was geschieht nun, wenn wir die Spannung zwischen Basis und Emitter größer als die Diffusionsspannung machen, also z.B. auf 0,7 Volt erhöhen? Vergleichen wir noch einmal mit der Kennlinie eines PN-Übergangs (Diode) für die Durchlaßrichtung: Sobald die angelegte Spannung in Durchlaßrichtung größer als die Diffusions- oder Schwellspannung ist, steigt der Strom steil an. Das trifft sinngemäß auch für unseren PN-Übergang Basis-Emitter zu; der Elektronenstrom zur Basis steigt stark an. Das Angebot an freien Elektronen in der Basisschicht wird sehr groß. Damit steigen aber auch der Strom zum Kollektor- und der Strom zum Basisanschluß stark.

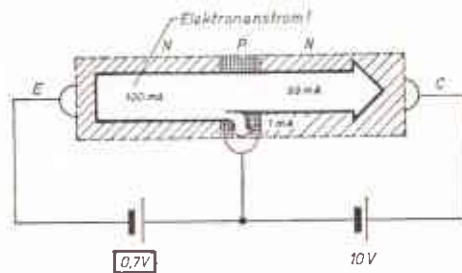


Abb. 4.6 — Stromverzweigung bei erhöhter Basis-Emitter-Spannung

Durch geringe Änderungen der Basis-Emitter-Spannung eines Transistors läßt sich also der Kollektorstrom stark beeinflussen. Da der Basisstrom jeweils sehr gering gegenüber dem Kollektorstrom ist, können wir ebenso feststellen, daß bei nur geringen Basisstromänderungen große Kollektorstromänderungen auftreten. Praktisch beträgt der Basisstrom nur etwa 1 % des Kollektorstroms.

Wir wollen uns merken: Bei einem Transistor werden durch geringe Änderungen der Basis-Emitter-Spannung bzw. des Basisstroms große Stromänderungen im Kollektor-Stromkreis hervorgerufen. Der Transistor hat somit eine Verstärkerwirkung.

In unseren Beispielen ändert sich infolge der Spannungsänderung von 0,5 Volt auf 0,7 Volt zwischen Basis und Emitter

- a) der Basisstrom von 0,01 mA auf 1 mA, also um 0,99 mA,
- b) der Kollektorstrom von 0,99 mA auf 99 mA, also um 98,01 mA.

Bilden wir das Verhältnis von Kollektorstromänderung ΔI_C zu Basisstromänderung ΔI_B , so erhalten wir den sogenannten Wechselstrom-Verstärkungsfaktor β . (Gilt für die Emitterschaltung.)

Für unsere Beispiele ergibt sich

$$\beta = \frac{\Delta I_C}{\Delta I_B} = \frac{98,01 \text{ mA}}{0,99 \text{ mA}} = 99 \approx 100$$

Wir haben festgestellt, daß sich am Kollektorstrom eines Transistors so lange nichts ändert, wie sich auch an der Basis-Emitter-Spannung bzw. am Basisstrom nichts ändert. Mit anderen Worten: Zu einem bestimmten Basisstrom gehört ein für den Transistortyp bestimmter Kollektorstrom.

Der Basisstrom wird bei allgemeinen Verstärkeranwendungen des Transistors mit Hilfe der Basis-Emitter-Spannung immer so eingestellt, daß Änderungen sowohl zu größeren als auch zu kleineren Stromstärken hin möglich sind, ohne den Transistor zu überlasten.

Dieser für den Transistortyp bestimmte Strom wird auch als **Basisruhestrom** und der zugehörige Kollektorstrom als **Kollektorruhestrom** bezeichnet. Beide liegen als zugeordnete Werte fest. Man bezeichnet diese Werte als Arbeitspunkte. Wie bereits im Abschn. 3 „Halbleiterdioden“ beschrieben, versteht man in der Elektronik unter dem Arbeitspunkt eines Bauelements ganz bestimmte, für die richtige Funktion des Bauelements festgelegte Gleichspannungs- oder Gleichstromwerte. Die günstigsten Arbeitspunkte werden für die Bauelemente üblicherweise vom Hersteller in Datenblättern angegeben.

Auch für den Arbeitspunkt eines Transistors ergibt sich ein ganz bestimmtes Verhältnis von Kollektorruhestrom I_C zu Basisruhestrom I_B .

Dieses Verhältnis wird als Gleichstromverstärkungsfaktor B bezeichnet und gilt ebenfalls für die Emitterschaltung des Transistors.

In unseren Beispielen liegt der Arbeitspunkt bei einer Basis-Emitter-Spannung von 0,6 Volt. Dabei betragen $I_C = 9,9 \text{ mA}$ und $I_B = 0,1 \text{ mA}$. Somit beträgt der

Gleichstromverstärkungsfaktor

$$B = \frac{I_C}{I_B} = \frac{9,9 \text{ mA}}{0,1 \text{ mA}} = 99$$

Vergleicht man nun den Gleichstromverstärkungsfaktor B mit dem Wechselstromverstärkungsfaktor β , so ergeben sich nahezu gleiche Werte. Der Einfachheit halber wird daher für einen Transistor in den Datenblättern oder Tabellen meistens der Gleichstromverstärkungsfaktor B angegeben. Er ist auch — wie wir noch durch Übungen bestätigt finden werden — mit verhältnismäßig einfachen Mitteln meßbar.

Nachdem wir die grundsätzliche Wirkungsweise eines Transistors anhand schematischer Abbildungen kennengelernt haben, wollen wir künftig für den Transistor das übliche Schaltsymbol verwenden. Folgende Symbole gelten für Transistoren:

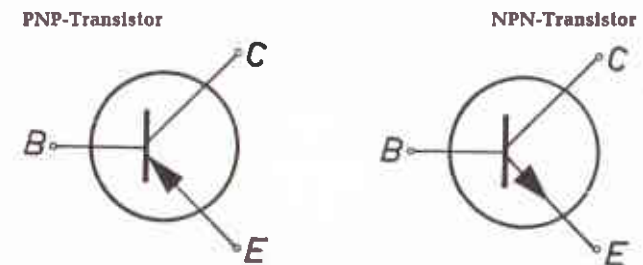


Abb. 4.7 — Schaltsymbole der Transistoren

Das Lesen von Schaltzeichen mit Transistoren wird dadurch erleichtert, daß die Pfeilspitze im Transistorsymbol immer in die Durchlaßrichtung des Stroms für den PN-Übergang Basis-Emitter weist (techn. Stromrichtung!). Zum besseren Verständnis der grundsätzlichen Funktion eines Transistors wurden die äußeren Spannungsquellen bisher immer direkt an Basis-Emitter bzw. an Basis-Kollektor geschaltet. Da aber die beiden Spannungsquellen in Reihe geschaltet sind, kann man sie auch anders an den Transistor legen:

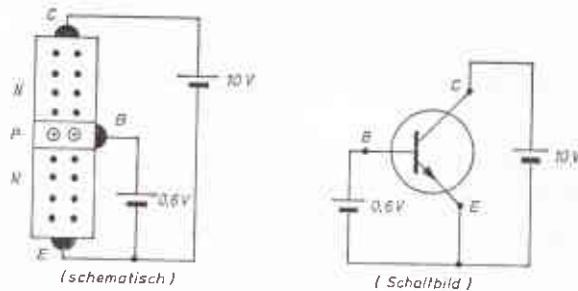


Abb. 4.8

An der Wirkungsweise des Transistors ändert sich dabei grundsätzlich nichts. Der Vorteil dieser Schaltung liegt u.a. darin, daß hier über die Basis-Emitter-Spannungsquelle nur der Basisstrom fließt. Somit wird nur eine geringe Steuerleistung benötigt. In der ursprünglichen Schaltung floß nämlich auch der Kollektorstrom über die Basis-Emitter-Spannungsquelle.

4.1.1. Bauformen und Anschlußarten der Transistoren

Um die Transistoren zu den in den nächsten Abschnitten vorgesehenen Messungen richtig anschließen zu können, hier erst einmal ein Überblick über Bauformen und Anschlußarten. Dazu muß allerdings von vornherein festgestellt werden, daß es unmöglich ist, sämtliche überhaupt vorkommenden Arten aufzuzeigen. Es gibt heute immerhin schätzungsweise einige 10 000 verschiedene Transistortypen! Wir müssen uns daher auf die Auswahl der gebräuchlichsten Arten beschränken.

Die Bezeichnung der Halbleiter-Bauelemente ist im Anhang (Abschn. 10) dieses Buches erläutert.

Das eigentliche Kristallsystem eines Kleinleistungstransistors hat nur Ausmaße in der Größenordnung um 100 bis 500 μm und wurde bisher überwiegend in Glas- oder Metallgehäuse eingebaut. In letzter Zeit werden Transistoren jedoch in zunehmendem Maße auch mit Kunststoffkapselung geliefert. Die kunststoffgekapselten Transistoren

sind aus fertigungstechnischen Gründen billiger als Transistoren mit Metallgehäuse. Kunststoffe leiten jedoch die Wärme erheblich schlechter als Metalle. Aus diesem Grunde findet man für Transistoren ein und desselben Typs in den Datenblättern bei Metallkapselung größere zulässige Verlustleistungen als bei Kunststoffkapselung. Näheres über Verlustleistung und alle Probleme, die mit der Wärmewirkung des Stroms im Innern des Transistors zusammenhängen, behandeln wir noch im Abschn. 4.4.5 „Arbeitspunktstabilisierung der Transistorverstärker“.

Beispiel:

Type	BC 107	$\cong$	BC 147
Gehäuse	Metall (TO-18)		Kunststoff (SOT-25)
zul. Verlustleistung	300 mW		220 mW.

Alle übrigen Daten sind für beide Typen gleich!

Eine Zusammenstellung gebräuchlicher Metallgehäuse-Typen ist im Abschn. 10 wiedergegeben.

Zur Bestimmung der Transistoranschlüsse müssen wir den Transistor von unten her — also von der Anschlußseite — betrachten. Zur Festlegung der Zählrichtung dient bei älteren Typen ein roter Farbpunkt an der Kollektor-Anschlußseite. Die neueren Metallgehäuse (Bezeichnung TO und folgende Nr.) besitzen häufig an der Anschlußseite eine kleine Blechfahne, die zur Orientierung dient. Von unten auf den Transistor gesehen, zählt man dann von der Fahne beginnend rechts herum, z.B. E—B—C.

Manchmal enthalten Transistoren einen vierten Anschlußdraht s , der mit dem Metallgehäuse verbunden ist. Dieser Anschluß dient zum Verbinden des Gehäuses mit dem Nullpotential (z.B. Chassis) einer Schaltung. Damit wird eine elektrostatische Abschirmung des Transistor Systems gegen Fremdfelder bewirkt (vor allem bei Hochfrequenz-Transistoren). Bei einigen Transistortypen, z.B. auch bei dem für unsere Versuche verwendeten Typ BC 107, ist der Kollektoranschluß mit dem Metallgehäuse verbunden. Leistungstransistoren besitzen aus wärmetechnischen Gründen ein sehr kompaktes Metallgehäuse, aus dem im allgemeinen nur zwei Anschlüsse als Stifte herausgeführt sind — nämlich für Emitter und Basis. Der Kollektor ist mit dem Metallgehäuse verbunden, da an ihm die weitaus stärkste Wärmeentwicklung im Transistor auftritt.

Da bei jedem Transistor Emitter- und Kollektorschicht grundsätzlich aus gleichartigem Halbleiterstoff bestehen (also entweder beides P-Leiter oder beides N-Leiter), können diese beiden Anschlüsse nicht ohne weiteres durch eine Messung bestimmt werden. Sie dürfen aber auf keinen Fall vertauscht werden. Da nämlich am Kollektor die stärkste Erwärmung im Betriebszustand auftritt, hat die Kollektorschicht größere Ausmaße als die Emitterschicht. Die größere Oberfläche der Kollektorschicht ermöglicht eine schnellere und bessere Wärmeableitung an die Umgebung.

Im Zweifelsfall muß daher immer auf die Hersteller-Datenlisten oder Tabellenbücher zurückgegriffen werden. Auf die meßtechnische Unterscheidungsmöglichkeit zwischen Emitter und Kollektor kommen wir im Abschn. 4.2 „Einfache Messungen an Transistoren“ zurück.

In den folgenden Abbildungen ist eine Auswahl gebräuchlicher Bauformen für Transistoren zusammengestellt. Die angegebenen Werte für die zulässige Verlustleistung sind als Richtwerte zu verstehen. Abweichungen sind sowohl nach oben als auch nach unten möglich.

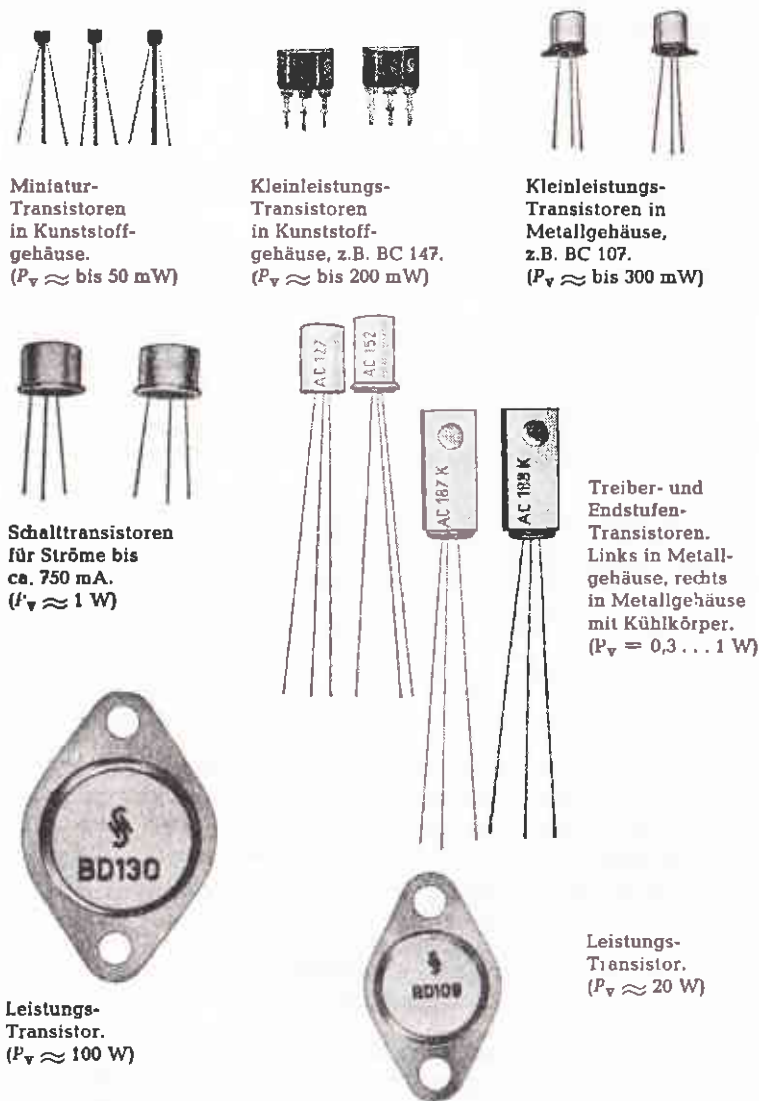


Abb. 4.9 — Auswahl gebräuchlicher Transistor-Bauformen

4.1.2. Die Bezeichnung und Zählrichtung der Ströme und Spannungen eines Transistors

Damit Datenblätter und -tabellen richtig gelesen werden können, sind hier die wichtigsten Bezeichnungen der Ströme und Spannungen sowie deren Zählrichtung (Vorzeichen) zusammengestellt.

4.1.2.1. Bezeichnung und Zählrichtung der Ströme

Die in den drei Anschlußdrähten eines Transistors im Betriebszustand fließenden Ströme werden so bezeichnet:

Kollektorstrom I_C (auch Kollektorruestrom),
 Basisstrom I_B (auch Basisruhestrom),
 Emittorstrom I_E (auch Emittorruestrom).

Bei der Zählrichtung der Ströme ist grundsätzlich für alle Transistoren festgelegt worden: Für alle Ströme, die in den Transistor hinein-fließen, gilt positive Zählrichtung, also auch positives Vorzeichen. Demnach muß vor Angaben über Ströme, die aus dem Transistor ab-fließen, ein negatives Vorzeichen gesetzt werden.

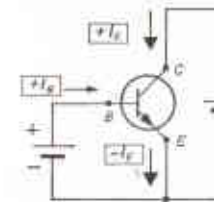


Abb. 4.10

Beispiel NPN-Transistor: Im Betriebszustand liegt sowohl an der Basis als auch am Kollektor positives Potential gegenüber dem Emittor. Basis- und Kollektorstrom fließen also in den Transistor hinein. Der Emittorstrom fließt aber ab und daher in entgegengesetzter Richtung zur festgelegten Zählweise. Somit erhalten I_C und I_B positives und I_E negatives Vorzeichen.

4.1.2.2. Bezeichnung und Zählrichtung der Spannungen

Die zum Betrieb am Transistor liegenden Spannungen werden in der Bezeichnung jeweils durch die beiden Anschlüsse des Transistors gekennzeichnet, zwischen denen die Spannung wirksam ist, also

Spannung zwischen Kollektor und Emittor: U_{CE} ,
 Spannung zwischen Kollektor und Basis: U_{CB} und
 Spannung zwischen Basis und Emittor: U_{BE} .

Als positive Zählrichtung gilt bei Spannungsangaben die Reihenfolge der Transistoranschlüsse, wobei immer der erste Anschluß positiv gegenüber dem zweiten sein soll. Liegt eine Spannung in der Polarität aber genau umgekehrt, so wird das durch ein negatives Vorzeichen ausgedrückt.

Beispiel PNP-Transistor: Zum Betrieb eines PNP-Transistors muß an seinem Kollektor und an seiner Basis negatives Potential gegenüber dem Emitter liegen. Daher gilt hier:

Kollektor-Emitter-Spannung: $-U_{CE}$ oder auch $U_{CE} = -\dots V$ (denn der Kollektor ist gegenüber dem Emitter negativ), Basis-Emitter-Spannung: $-U_{BE}$ oder auch $U_{BE} = -\dots V$ (denn die Basis ist gegenüber dem Emitter negativ) und Kollektor-Basis-Spannung $-U_{CB}$ bzw. $U_{CB} = -\dots V$.

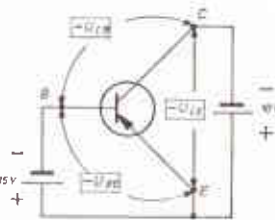


Abb. 4.11

Wichtig ist, daß diese Spannungen jeweils direkt zwischen den Transistoranschlüssen zu messen sind (und nicht etwa an Vorwiderständen)! Ist in Datenblättern oder -tabellen die Schichtfolge NPN oder PNP nicht ausdrücklich angegeben, so kann sie aus den Vorzeichen der Strom- und Spannungswerte abgeleitet werden.

Beispiel BC 179:

Grenzdaten: $-U_{CEmax} = 20 V$ oder $U_{CEmax} = -20 V$,
 $-I_{Cmax} = 100 mA$ oder $I_{Cmax} = -100 mA$.

Es muß sich um einen PNP-Transistor handeln, da der Kollektor gegenüber dem Emitter negatives Potential erhält und der Kollektorstrom vom Transistor abfließt.

4.2. Einfache Messungen an Transistoren

Ähnlich wie bei den Dioden wollen wir Art und Funktionsfähigkeit von Transistoren mit Hilfe einfacher Messungen mit einem Vielfachinstrument prüfen. Diese einfachen Messungen lassen zwar keine Rückschlüsse auf die genauen Betriebsdaten der Bauelemente zu. Sie ermöglichen es uns jedoch, die Schichtfolge eines Transistors, seine Funktionsfähigkeit bzw. grobe Fehler verhältnismäßig einfach zu bestimmen.

4.2.1. Ermittlung der Schichtfolge (PNP oder NPN) eines Transistors und Prüfung seiner Funktionsfähigkeit

Voraussetzung für eine einfache Prüfung ist, daß wir die richtige Reihenfolge der Transistoranschlüsse oder zumindest einen der drei Anschlüsse (C oder E) kennen. Unsere Prüfung wird durch die Tatsache ermöglicht, daß ein Transistor grundsätzlich zwei PN-Übergänge enthält und daher für einfache Messungen aus zwei Dioden zusammengesetzt gedacht werden kann.

Wir möchten jedoch darauf hinweisen, daß das häufig in der Literatur zu findende Ersatzschaltbild eines Transistors — zwei gegeneinandergeschaltete Dioden — nicht die Funktion als Verstärker erkennen läßt; denn im Transistor stellt immer die Basis-schicht eine gemeinsame Schicht für beide PN-Übergänge dar. Man kann also auf keinen Fall einen Transistor herstellen, indem man zwei einfache Dioden in entgegengesetzter Durchlaßrichtung hintereinanderschaltet! Für unsere einfachen Messungen ist das Dioden-Ersatzschaltbild jedoch sehr einfach und anschaulich.

Für die beiden grundsätzlich möglichen Arten von Transistoren ergeben sich also folgende Ersatzschaltbilder:

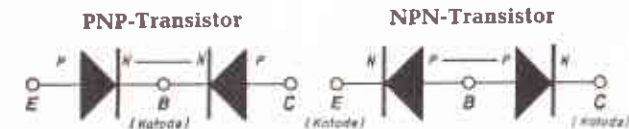


Abb. 4.12 — Ersatzschaltbilder für Transistoren

Wir können somit für beide Arten die Strecke E—B und die Strecke B—C als je eine Diode auffassen. Die Funktionsprüfung einer Diode ist bereits bekannt; unsere Erfahrungen können jetzt sinngemäß für die Prüfung von Transistoren angewandt werden. Für unsere erste Messung verwenden wir den Transistor BC 107 (oder BC 108 oder BC 109). Nach dem Datenblatt gilt die untenstehende Anschlußreihenfolge.

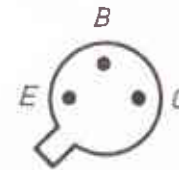


Abb. 4.13 — BC 107 (Anschlußseite)

Von der Metallfahne am Gehäuse angefangen liegen also nacheinander Emitter-, Basis- und Kollektoranschluß rechtsherum.

Eine Diodenprüfung wird durch eine Widerstandsmessung mit einem Ohmmeter vorgenommen und der Meßbereich des Ohmmeters oder Vielfachinstruments auf $\times 1k$ geschaltet. Die erste Messung gilt der Basis-Emitter-Diode. Für diese Diodenstrecke ermitteln wir in üblicher Weise Durchlaß- und Sperrrichtung sowie Durchlaß- und Sperrwiderstand. Da es sich in unserem Fall um einen Silizium-Transistor handelt, ist der Sperrwiderstand mit einem einfachen Instrument nicht mehr meßbar (er hat einen Wert von rd. 10 Megohm). Wir können aber eindeutig Durchlaß- und Sperrrichtung unterscheiden und damit den Katodenanschluß dieser Diodenstrecke bestimmen. Die Messungen führen zu etwa folgendem Ergebnis:

Sperrwiderstand:	Anzeige ∞ ,
Durchlaßwiderstand:	ca. 10 bis 15 Kiloohm,
Katode:	Emitter,
folglich Durchlaßrichtung: (techn. Strom!)	Basis-Emitter.

Ergibt sich im Sperrbereich ein deutlicher Ausschlag des Ohmmeterzeigers, dann ist die Basis-Emitter-Diode nicht einwandfrei. Als nächstes folgt die gleiche Prüfung zwischen Basis- und Kollektoranschluß. Die Messungen ergeben

Sperrwiderstand:	Anzeige ∞ ,
Durchlaßwiderstand:	ca. 8 bis 12 Kiloohm,
Katode:	Kollektor,
folglich Durchlaßrichtung: (techn. Strom!)	Basis-Kollektor.

Auch hier gilt: Ist die Anzeige des Ohmmeters im Sperrbereich deutlich kleiner als ∞ , so ist diese Diodenstrecke nicht einwandfrei. Für unseren Transistor ergibt sich jetzt folgendes Ersatzschaltbild:



Abb. 4.14

Es handelt sich also um einen NPN-Transistor.

Für die Ermittlung der Schichtfolge an Germanium-Transistoren gelten die gleichen Messungen. Zu beachten ist aber dabei, daß sich bei Germanium-PN-Übergängen im allgemeinen nicht so große Sperrwiderstände ergeben wie bei Silizium-PN-Übergängen. Bei den obengenannten Typen muß sich die Schichtfolge PNP ergeben. Für die Messungen dieser Art gilt das gleiche wie für die einfachen Messungen an Dioden: **Das Verhältnis von Sperrwiderstand zu Durchlaßwiderstand sollte bei jedem der beiden PN-Übergänge eines Transistors etwa 100 : 1 oder größer sein.** Bei Kleinleistungs- und vor allem Siliziumtransistoren beträgt das Verhältnis mehr als 100. Dagegen kann es bei Leistungstransistoren auch noch darunter liegen.

4.2.1.1. Ermittlung der Anschlüsse bei unbekannter innerer Beschaltung

Ist kein Anschluß des zu prüfenden Transistors bekannt, so sind folgende Prüfungen vorzunehmen:

- Man ermittelt durch Messungen den **Basisanschluß**. Dazu wird jeweils eine Meßschnur des Ohmmeters zunächst fest an einen beliebigen Transistoranschluß gelegt und der Widerstand zu den anderen Transistoranschlüssen gemessen. Diese Messung wird nach Wechseln des ersten Anschlußdrahts bzw. durch Vertauschen der Meßschnüre so oft wiederholt, bis sich vom festen Anschluß zu den übrigen beiden ein niedriger Widerstandswert ergibt. In diesem Fall ist bei einem funktionsfähigen Transistor der erste, fest mit einer Meßschnur verbundene Anschlußdraht der Basisanschluß.
- Die **Funktionsfähigkeit der beiden PN-Übergänge** wird dadurch geprüft, daß von dem gefundenen Basisanschluß aus auch die Sperrwiderstände zu den beiden anderen Transistoranschlüssen gemessen werden. Ergibt sich auch nach Umpolen der Meßschnüre ein niedriger Widerstandswert, so ist der Transistor nicht in Ordnung.
- Die **Funktionsfähigkeit des Transistors** wird durch eine Widerstandsmessung zwischen Kollektor und Emitter überprüft. Die Polarität spielt dabei eine untergeordnete Rolle. Der gemessene Widerstand muß für beide Richtungen einen hohen Wert ergeben. (Er liegt aber im allgemeinen etwas unter dem Sperrwiderstand eines einzelnen PN-Übergangs.)
- Den **Kollektoranschluß** kann man — wenn auch nicht immer sehr zuverlässig — durch **genaues** Messen der Durchlaßwiderstände für die beiden PN-Übergänge bestimmen, und zwar von der Basis aus. Dabei ergibt sich fast immer, daß einer der gemessenen Durchlaßwiderstände geringfügig kleiner ist als der andere. **Der etwas kleinere Durchlaßwiderstand gehört zum Basis-Kollektor-PN-Übergang.** Das hängt damit zusammen, daß die Kollektorschicht eines Transistors einen größeren Querschnitt hat als die Emitterschicht. Da der Basisanschluß bereits gefunden wurde, liegt jetzt auch der Kollektoranschluß fest. Demnach muß der dritte Anschluß der Emitteranschluß sein. (Vergleichen Sie mit den Meßergebnissen am Transistor BC 107I)
- Die **Schichtfolge** des Transistors wird wie bereits beschrieben ermittelt.

Nach dieser Methode lassen sich für nahezu alle Transistoren die drei Anschlüsse sowie die Schichtfolge ermitteln. Es ist bei den entsprechenden Messungen aber immer wieder festzustellen, daß der Unterschied zwischen den Durchlaßwiderständen B—E und B—C meistens sehr gering ist. Daher ist eine peinlich genaue Ablesung der beiden Werte unerlässlich!

Es empfiehlt sich, jeden Transistor durchzuprüfen, bevor er für weitere Messungen oder Schaltungen verwendet wird. Dadurch vermeidet man eine oft umständliche und zeitraubende Fehlersuche.

4.2.2. Messung des Gleichstromverstärkungsfaktors B

Nachdem wir in der Lage sind, die drei Anschlüsse eines Transistors zu bestimmen, können wir ihn jetzt auch richtig für die folgenden Messungen und Schaltungen anschließen. Wie bereits angegeben, ergibt sich der Gleichstromverstärkungsfaktor B eines Transistors aus dem Verhältnis von Kollektorruhestrom zu Basisruhestrom:

$$B = \frac{I_C}{I_B}$$

Zur Messung des Gleichstromverstärkungsfaktors werden noch die Daten für den richtigen Arbeitspunkt benötigt.

Zum Gleichstromverstärkungsfaktor eines Transistors ist noch folgendes zu bemerken: In den Datenblättern wird entweder für einen angegebenen Arbeitspunkt ein ganz bestimmter Wert für B angegeben oder aber ein Bereich, innerhalb dessen der Verstärkungsfaktor liegen soll. Dies ist darauf zurückzuführen, daß Transistoren in der Fertigung erheblichen Streuungen unterliegen, so daß es unmöglich ist, Transistoren gleichen Typs mit genau gleichen Verstärkungsfaktoren herzustellen. Abweichungen um -50% bis zu $+100\%$ vom angegebenen Mittelwert sind durchaus möglich! Die angegebenen Festwerte für B sind also nur als Mittelwerte zu verstehen.

Der Verstärkungsfaktor hängt weiter, wie wir noch durch Messungen bestätigt finden werden, vom jeweiligen Arbeitspunkt ab. Will man sich also ein Bild vom tatsächlichen Verstärkungsfaktor eines bestimmten Transistors machen, so bleibt nur die Messung als Mittel übrig.

Da sich einerseits die Klemmenspannung einer Batterie während längerer Belastung ändert und andererseits Halbleiter-Bauelemente im Betrieb erwärmen, sind alle Messungen so kurzzeitig wie möglich vorzunehmen. Zur Vermeidung von Fehlern wird daher die Verwendung einer Taste zum kurzzeitigen Einschalten empfohlen.

Für unsere Messungen wollen wir wieder den Transistor BC 107 verwenden. Folgende Schaltung ist zusammenzustellen:

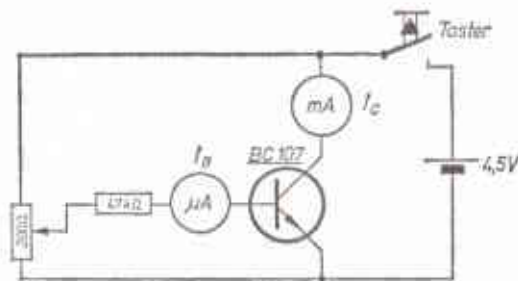


Abb. 4.15 — Messung des Gleichstromverstärkungsfaktors B

Für den Transistor BC 107 findet man im Datenblatt unter anderem folgende Angaben:

P_V	=	300 mW,
U_{CE}	=	5 V,
I_C	=	2 mA,
B	=	180.

BC 107, Kenndaten im Arbeitspunkt:

Ohne einen merklichen Fehler zu machen, kann der Einfachheit halber als Spannungsquelle eine Flachbatterie 4,5 V verwendet werden. Diese Spannungsquelle können wir sowohl für U_{CE} als auch für U_{BE} ausnutzen. Um eine Überbelastung des Transistors zu vermeiden, wird in den Basisstromkreis ein Schutzwiderstand geschaltet.

Berechnung des Schutzwiderstands:

Aus den gegebenen Werten für I_C und B läßt sich der zu erwartende Basisstrom I_B berechnen:

$$I_B = \frac{I_C}{B} = \frac{2 \text{ mA}}{180} = 0,011 \text{ mA} = 11,1 \mu\text{A}.$$

Da die Basis-Emitter-Strecke in Durchlaßrichtung betrieben wird, ist ihr Widerstand vernachlässigbar klein. Als Spannungsquelle dient die 4,5-V-Batterie. Im Stromkreis Basis-Emitter muß demnach ein Gesamtwiderstand von

$$R = \frac{U}{I_B} = \frac{4,5 \text{ V}}{11,1 \mu\text{A}} = \frac{4,5 \text{ V}}{11,1 \cdot 10^{-6} \text{ A}} = 0,405 \cdot 10^{-6} \Omega = 405 \text{ k}\Omega$$

liegen. Da sich jedoch die Verlustleistung für den gegebenen Arbeitspunkt nur zu

$$P_V = U_{CE} \cdot I_C = 4,5 \text{ V} \cdot 2 \cdot 10^{-3} \text{ A} = 9 \cdot 10^{-3} \text{ W} = 9 \text{ mW}$$

gegenüber 300 mW zulässiger Verlustleistung ergibt, kann der Schutzwiderstand ohne weiteres kleiner gewählt werden. Für unseren Aufbau verwenden wir einen 50-Kilo-ohm-Widerstand (nächster Normwert 47 kΩ).

Die Messung geht folgendermaßen vor sich: Ein Meßinstrument (Meßbereich 2,5 mA) zeigt den Kollektorstrom I_C an. Bei eingeschalteter Batterie (Taste benutzen!) wird jetzt mit Hilfe des 200-Ohm-Potentiometers die Basis-Emitter-Spannung und damit der Basisstrom I_B so eingeregelt, daß sich ein Kollektorstrom von genau 2 mA ergibt. Damit ist der vorgesehene Arbeitspunkt erreicht. Für diesen Arbeitspunkt muß nun auch der Basisstrom I_B ermittelt werden. Er wird mit einem Strommesser (50 μA -Meßbereich) gemessen.

Im Verhältnis zum gesamten Stromkreiswiderstand ist der Innenwiderstand des Meßinstrumentes so klein, daß sich praktisch kein Fehler durch das Einschalten des Instruments ergibt.

Für einen durchgemessenen Transistor BC 107 ergaben sich folgende Meßwerte:

$$I_C = 2 \text{ mA}; \quad I_B = 12 \text{ } \mu\text{A}.$$

Demnach beträgt der Gleichstromverstärkungsfaktor:

$$\beta = \frac{I_C}{I_B} = \frac{2 \text{ mA}}{12 \text{ } \mu\text{A}} = \frac{2 \text{ mA}}{0,012 \text{ mA}} = \underline{\underline{167}}.$$

An diesem Meßbeispiel ist bereits zu erkennen, daß zwischen der Angabe im Datenblatt und dem tatsächlichen Wert für einen Transistor dieser Type Abweichungen auftreten.

Die angegebene Meßschaltung kann zur Messung von Gleichstromverstärkungsfaktoren für beliebige NPN-Transistoren angewendet werden. Dabei ist darauf zu achten, daß der Schutzwiderstand R_V dem jeweiligen Transistortyp in etwa entspricht. Ebenso kann aber auch β für PNP-Transistoren bestimmt werden. Man braucht dazu lediglich die Batterie umzupolen.

4.3. Kennlinien des Transistors (Emitterschaltung)

Wir haben uns anhand einfacher Messungen mit dem Transistor grundsätzlich vertraut gemacht und wollen jetzt darangehen, ihn meßtechnisch genauer zu untersuchen. Die bisher bei Messungen gemachten Erfahrungen sind dabei sehr nützlich. Um einen Transistor und sein Verhalten vollkommen beurteilen zu können, ist die Aufnahme von Kennlinien notwendig. Die Kennlinie einer Gleichrichter-Diode und einer Z-Diode haben bereits erkennen lassen, wie einfach und anschaulich das Verhalten eines Halbleiter-Bauelements wird, wenn man seine Kennlinie vor Augen hat. Das Verhalten eines Transistors wird im wesentlichen durch drei verschiedene Kennlinienarten veranschaulicht, nämlich die **Eingangskennlinie**, die **Ausgangskennlinie** und die **Steuerkennlinien**.

Die für die Messungen gewählten Spannungs- und Stromwerte sind bewußt so festgelegt worden, daß sie einigermaßen in die Meßbereiche eines einfachen Vielfach-Meßinstruments fallen. Wenn diese Werte auch geringfügig außerhalb der vom Hersteller empfohlenen Arbeitspunkte liegen, so geben die aufgenommenen Kennlinien trotzdem ein vollständiges Bild von der Arbeitsweise des untersuchten Transistors.

4.3.1. Die Eingangskennlinie eines Transistors

Die **Eingangskennlinie** stellt die **Abhängigkeit des Eingangsstroms I_B von der Höhe der Basis-Emitter-Spannung U_{BE}** dar, wobei eine konstante Spannung U_{CE} zwischen Kollektor und Emitter anliegt. Die Meßschaltung sieht folgendermaßen aus:

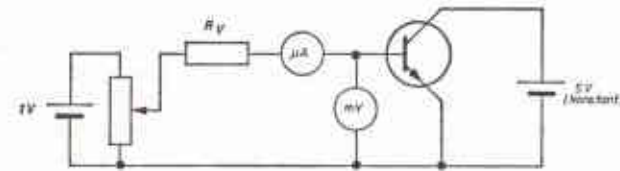


Abb. 4.16 — Meßschaltung zur Aufnahme der Eingangskennlinie

Da zur Aufnahme der Eingangskennlinie sehr kleine Ströme und sehr kleine Spannungen gemessen werden müssen, ist das Ergebnis nur dann brauchbar, wenn die Spannungen mit einem sehr hochohmigen Transistor- oder Röhrenvoltmeter gemessen werden. Die Basis-Emitter-Strecke in einem Transistor stellt einen PN-Übergang dar. Deshalb zeigt die Eingangskennlinie auch einen ähnlichen Verlauf wie eine Diodenkennlinie im Durchlaßbereich.

Wegen der Schwierigkeit, eine Eingangskennlinie mit einfachen Zeigerinstrumenten aufzunehmen, ist die folgende Kennlinie nach Herstellerunterlagen gezeichnet worden.

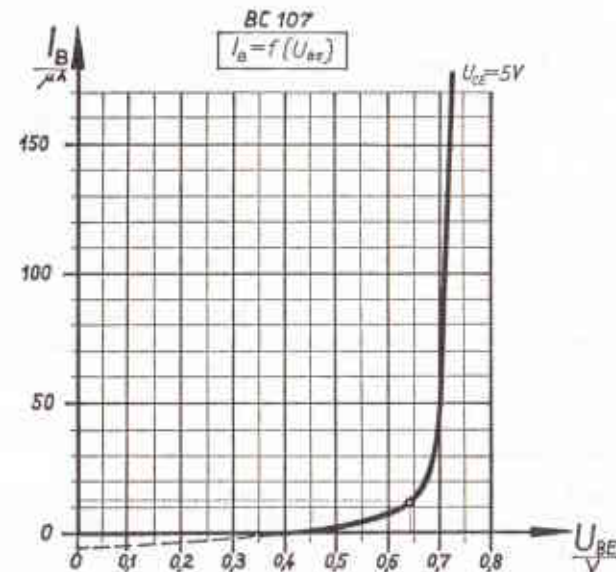


Abb. 4.17 — Eingangskennlinie eines Silizium-Transistors (BC 107)

Sie läßt erkennen, daß unterhalb der Diffusionsspannung des PN-Übergangs B—E verhältnismäßig große Änderungen der Basis-Emitter-Spannung nur geringe Basisstromänderungen bewirken. Erst wenn die Basis-Emitter-Spannung gleich oder größer als die Diffusionsspannung ist, machen sich Spannungsänderungen durch größere Stromänderungen bemerkbar. Zur Erzielung einer ausreichenden Verstärkerwirkung des Transistors muß also die Basis-Emitter-Spannung mindestens die Höhe der Diffusionsspannung haben. Überschläglich gelten folgende Spannungswerte:

Germanium-Transistoren $U_{BE} \approx 0,2 \dots 0,3 \text{ V}$,

Silizium-Transistoren $U_{BE} \approx 0,6 \dots 0,7 \text{ V}$.

Da durch die Höhe der Basis-Emitter-Spannung der Basisstrom und damit auch der Arbeitspunkt festliegt, gilt für die richtige Größe folgendes: Für Vorverstärker ist am Eingang nur mit kleinen Spannungsänderungen zu rechnen. In Vorverstärkern ist ein Rauschen unerwünscht; denn es würde ja von den nachfolgenden Stufen mit dem Signal weiterverstärkt. Da das Rauschen eines Transistorverstärkers mit höherem Arbeitspunkt steigt, wird die Basis-Emitter-Spannung für Vorverstärker verhältnismäßig niedrig gewählt (etwa gleich der Diffusionsspannung). Die nachfolgenden Verstärkerstufen sollen das Signal weiterverstärken. Sie erhalten also am Eingang bereits größere Spannungsänderungen. Ihr Arbeitspunkt muß daher durch eine entsprechend höhere Basis-Emitter-Spannung höhergelegt werden.

Aus der Eingangskennlinie kann man den statischen (Gleichstrom-) und den dynamischen (Wechselstrom-)Eingangswiderstand eines Transistors bestimmen. Die Höhe der konstanten Kollektor-Emitter-Spannung U_{CE} hat übrigens auf den Verlauf der Eingangskennlinie nur einen sehr unwesentlichen Einfluß, so daß man hier von der Aufnahme einer Kennlinienschar (etwa wie bei den Ausgangskennlinien) absehen kann.

Kennlinien stellen im mathematischen Sinn sogenannte Funktionen dar. Unter einer Funktion versteht man die Abhängigkeit einer Größe von einer anderen veränderlichen Größe. Da solche Abhängigkeiten — sprich Funktionen — in der Physik und in der Technik sehr häufig untersucht werden, wollen wir auch die dafür übliche Schreibweise kennenlernen.

Unsere Eingangskennlinie stellt als Funktion die Abhängigkeit des Basisstroms I_B von der Basis-Emitter-Spannung U_{BE} dar. Man schreibt daher $I_B = f(U_{BE})$ und liest:

„ I_B ist eine Funktion von U_{BE} “ oder einfach

„ I_B in Abhängigkeit von U_{BE} “.

Die während der Messungen **konstant bleibende Größe** — hier die Spannung U_{CE} — wird als **Parameter** bezeichnet.

Die Eingangskennlinie $I_B = f(U_{BE})$ eines Transistors läßt das Verhalten des PN-Übergangs Basis-Emitter bei veränderlicher Basis-Emitter-Spannung erkennen. Aus der Eingangskennlinie können der Eingangswiderstand und die für einen bestimmten Basisstrom (Arbeitspunkt) erforderliche Basis-Emitter-Spannung bestimmt werden.

Beispiel:

Damit im Arbeitspunkt bei $U_{CE} = 5 \text{ V}$ $I_C = 2 \text{ mA}$ über unseren Transistor BC 107 fließen, ist ein Basisstrom $I_B = 12 \mu\text{A}$ einzustellen (siehe Messung des Gleichstromverstärkungsfaktors). Aus der Eingangskennlinie können wir zu $I_B = 12 \mu\text{A}$ eine Spannung von $U_{BE} \approx 0,64 \text{ V}$ ablesen. Der gewünschte Arbeitspunkt liegt also auch dadurch fest, daß wir $0,64 \text{ V}$ zwischen Basis und Emitter anlegen (vgl. hierzu Abb. 4.17).

4.3.2. Die Ausgangskennlinien eines Transistors

Die Ausgangskennlinien stellen die Abhängigkeit des Kollektorstroms von der Spannung zwischen Kollektor und Emitter dar. Dabei wird entweder die Basis-Emitter-Spannung oder der Basisstrom konstant gehalten (Parameter). Da der Kollektorstrom jedoch auch vom Basisstrom (und dieser von der Basis-Emitter-Spannung) abhängt, nimmt man mehrere solcher Ausgangskennlinien auf, wobei zu jeder Kennlinie ein ganz bestimmter Wert der Basis-Emitter-Spannung oder des Basisstroms gehört (Kennlinienschar). Wegen der Schwierigkeit, die Basis-Emitter-Spannung auf einem einfachen Meßinstrument genau abzulesen, wollen wir die Kennlinien für jeweils konstanten Basisstrom aufnehmen und müssen dazu die folgende Schaltung sehr sorgfältig aufbauen.

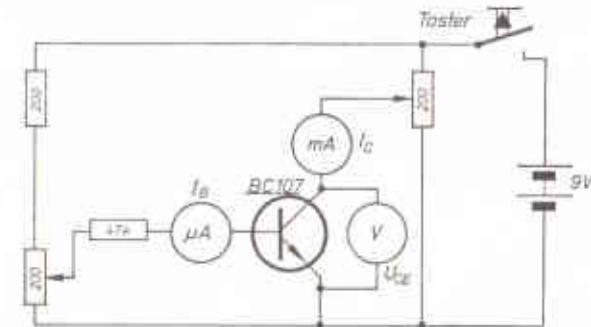


Abb. 4.18 — Meßschaltung zur Aufnahme der Ausgangskennlinien

Bei den Messungen ist zu beachten:

- Die Kollektorspannung wird zunächst in kleineren Stufen geändert, da die Kennlinien für kleine Spannungen noch sehr gekrümmt sind. Für den nahezu linearen Kennlinienbereich reichen größere Spannungsstufen aus.
- Alle drei Anschlüsse eines Transistors sind galvanisch untereinander verbunden. Infolgedessen ändern sich u.U. alle anderen Werte, wenn man eine Größe verändert. Man muß also

bei jeder eingestellten Stufe der Kollektorspannung unbedingt auch den richtigen Wert des Basisstroms überprüfen und ggf. nachregeln!

- c) Unter keinen Umständen darf die Meßschaltung zu lange an der Batterie liegen, da sich sonst durch Absinken der Batteriespannung oder Erwärmung des Transistors Meßfehler ergeben. Taste benutzen!
- d) Auf das Einschalten von Schutzwiderständen ist bereits hingewiesen worden. Das gilt ganz besonders für den Basisstromkreis.

Die Meßergebnisse tragen wir in vorbereitete Tabellen ein und zeichnen anschließend die Kennlinien.

Infolge kleiner Meßungenauigkeiten liegen nicht alle Meßpunkte genau auf den Kennlinien. Beim Ausziehen der Kennlinien gleicht man diese Abweichungen aus, verbindet also nicht alle Punkte pedantisch genau.

Zum Vergleich sind in Abb. 4.19 die Ausgangskennlinien für einen Transistor BC 107 dargestellt.

Die Kennlinien gelten für eine größere Typenzahl von Transistoren (wie aus der Angabe im Kopf hervorgeht). Dabei unterscheiden sich die vergleichbaren Typen nur durch die Art des Gehäuses und die zulässige Verlustleistung.

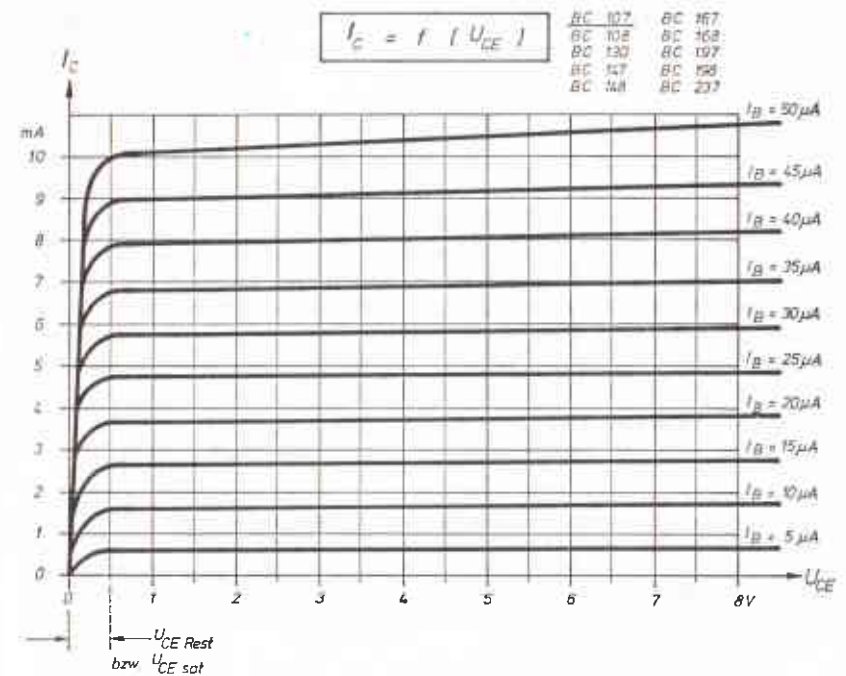


Abb. 4.19 — Ausgangskennlinien für einen Transistor BC 107

Welche Folgerungen lassen sich aus dem Kennlinienverlauf ziehen?

Der Kollektorstrom eines Transistors steigt erst merklich an, wenn die Kollektorspannung U_{CE} etwa die Höhe der Diffusionsspannung Basis-Emitter erreicht hat. Das bestätigt unsere Anschauung von der Arbeitsweise eines Transistors: Der Kollektorstrom ist größer als der Basisstrom, wenn am Kollektor ein höheres Potential liegt als an der Basis. **Damit also ein brauchbarer Kollektorstrom über den Transistor fließt, muß eine bestimmte Mindestspannung $U_{CE \text{ Rest}}$ angelegt werden, die auch vielfach als $U_{CE \text{ sat}}$ (Kollektor-Sättigungsspannung) bezeichnet wird.** Die von den Herstellern in den Datenblättern dargestellten Ausgangskennlinien eines Transistors beginnen daher häufig erst bei dieser **Kollektor-Emitter-Restspannung**. Die Beachtung der Restspannung $U_{CE \text{ Rest}}$ ist für die Auslegung von Schaltungen wichtig, bei denen nur mit kleinen Batteriespannungen gearbeitet werden soll. Die Spannung U_{CE} darf am Transistor im Betrieb auf keinen Fall unter den Wert von $U_{CE \text{ sat}}$ sinken, da er dann nicht mehr einwandfrei arbeitet.

Die Restspannung ist für die Arbeitsweise eines Transistors als Schalter von besonderer Bedeutung; sie muß für Schalttransistoren sehr klein sein.

Interessant für die Beurteilung der Arbeitsweise des Transistors ist die Tatsache, daß der Kollektorstrom mit höher werdender Kollektor-Emitter-Spannung U_{CE} oberhalb von $U_{CE_{Rest}}$ nur noch wenig zunimmt. Oberhalb $U_{CE_{Rest}}$ zeigt sich eine Art Sättigung im Kennlinienverlauf ähnlich wie bei einer Hystereseschleife. Der Grund ist einfach: Der Kollektor kann nicht mehr freie Ladungsträger absaugen als in die Basis vordringen. Die Zahl der in die Basis eindringenden freien Ladungsträger hängt aber ausschließlich von der Basis-Emitter-Spannung und somit vom Basisstrom ab. Da der Basisstrom für jede Kennlinie fest eingestellt wird, liegt auch die Zahl der freien Ladungsträger in der Basiszone fest. Der flache Verlauf der Ausgangskennlinien weist darauf hin, daß der **dynamische Ausgangswiderstand groß ist**.

Daß der Kollektorstrom dennoch mit steigender Spannung geringfügig zunimmt, ist damit zu erklären, daß mit zunehmender Gesamtspannung zwischen Kollektor und Emitter auch die Teilspannung am Basis-Emitter-Widerstand größer wird. Der Transistor stellt ja mit seinen beiden Teilwiderständen R_{BE} und R_{BC} einen Spannungsteiler dar. Infolge der etwas höheren Basis-Emitter-Spannung steigt somit auch die Zahl der in die Basis vordringenden freien Ladungsträger und damit der Kollektorstrom geringfügig an.

Die Ausgangskennlinien eines Transistors veranschaulichen sein elektrisches Verhalten zwischen Kollektor und Emitter, z.B. sein Widerstandsverhalten. Anhand des Ausgangskennlinienfeldes kann der günstigste Wert für den Lastwiderstand im Kollektorstromkreis ermittelt werden.

4.3.3. Die Steuerkennlinien eines Transistors

Wie die Bezeichnung schon erkennen läßt, veranschaulichen die Steuerkennlinien die Steuerwirkung, also die Verstärkerwirkung eines Transistors. Man kann für einen Transistor zwei Arten von Steuerkennlinien aufnehmen:

- Abhängigkeit des Kollektorstroms von der Basis-Emitter-Spannung $I_C = f(U_{BE})$ (Spannungs-Steuerkennlinie) oder
- Abhängigkeit des Kollektorstroms von der Stärke des Basisstroms $I_C = f(I_B)$ (Strom-Steuerkennlinie).

Bei der Aufnahme beider Kennlinienarten wird die Kollektor-Emitter-Spannung U_{CE} jeweils fest eingestellt (Parameter). Ihre Höhe hat — ähnlich wie bei der Eingangskennlinie — nur einen geringen Einfluß auf den Kennlinienverlauf. Wegen der Schwierigkeit, kleine Spannungen

im Millivoltbereich mit einfachen Instrumenten zu messen, ist die folgende Spannungs-Steuerkennlinie nach Herstellerunterlagen gezeichnet worden.

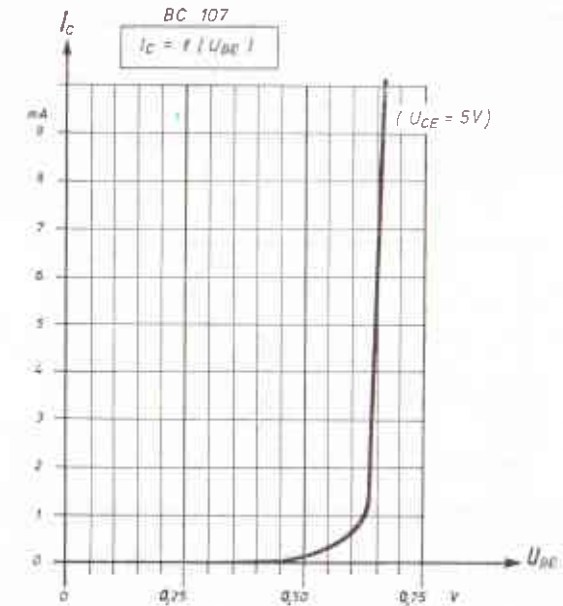


Abb. 4.20 — Spannungs-Steuerkennlinie BC 107

Auf den ersten Blick zeigt sich hier nahezu eine Übereinstimmung mit der Eingangskennlinie. Der einzige Unterschied besteht in der senkrechten Stromachse, auf der hier der Kollektorstrom I_C abgetragen wird. Der gleiche Verlauf wie bei der Eingangskennlinie ergibt sich aus einfachem Zusammenhang:

$$\text{Da } B = \frac{I_C}{I_B}, \text{ ist } I_C = B \cdot I_B.$$

Der Kollektorstrom ist also um den Gleichstromverstärkungsfaktor größer als der Basisstrom. Wir können, ohne einen groben Fehler zu machen, die Steuerkennlinie $I_C = f(U_{BE})$ aus der Eingangskennlinie des betreffenden Transistors bestimmen, wenn der Gleichstromverstärkungsfaktor bekannt ist.

Die Steuerkennlinie $I_C = f(U_{BE})$ verläuft stark gekrümmt. Daraus folgt: Wird die Basis-Emitter-Spannung U_{BE} gleichmäßig verändert, so ändert sich der Kollektorstrom nicht gleichmäßig (**nichtlinear**). Als Wechselspannungsverstärker — z.B. für Tonfrequenzen — wird also der Transistor um so stärkere Verzerrungen verursachen, je größer die Spannungsschwankungen am Eingang des Transistors sind. Sollen trotzdem große Spannungsschwankungen möglichst verzerrungsfrei verstärkt werden, so muß man auf dem steil ansteigenden Teil der

Kennlinie arbeiten, da dieser geradliniger verläuft. Das bedeutet, daß der Arbeitspunkt um so höher gelegt werden muß, je größer die Steuerwechselspannung ist.

Etwas anders sieht es nun mit der Strom-Steuerkennlinie $I_C = f(I_B)$ aus. Sie läßt erkennen, wie sich der Kollektorstrom ändert, wenn der Basisstrom verändert wird. Eine solche Kennlinie wollen wir für unseren Transistor BC 107 aufnehmen. Dazu muß die folgende Schaltung aufgebaut werden (sie gleicht der Meßschaltung zur Bestimmung des Gleichstromverstärkungsfaktors).

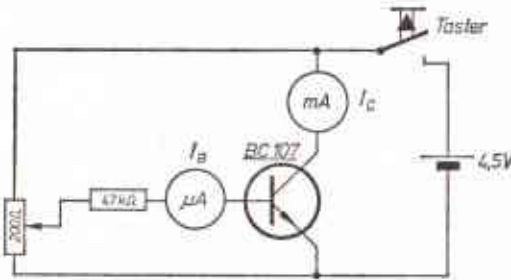


Abb. 4.21 — Schaltung zur Aufnahme der Strom-Steuerkennlinie $I_C = f(I_B)$

Die Spannung U_{CE} zwischen Kollektor und Emitter entnehmen wir einer Flachbatterie 4,5 Volt. Sie bleibt während der Kennlinienaufnahme unverändert. Mit Hilfe des Potentiometers wird jetzt der Basisstrom in Stufen vergrößert und zu jeder Stufe der Kollektorstrom gemessen. Unsere Ergebnisse tragen wir wieder in eine Tabelle ein. Danach zeichnen wir die Kennlinie.

Die Steuerkennlinie läßt erkennen, daß der **Verlauf nahezu linear ist**. Durch Änderung des Basisstroms ändert sich der Kollektorstrom praktisch im gleichen Verhältnis. Soll also eine möglichst verzerrungsfreie Verstärkung erreicht werden, so ist der Transistor mit veränderlichem **Strom** (Wechselstrom) anzusteuern. Man spricht dabei auch von der **Stromsteuerung** eines Transistors.

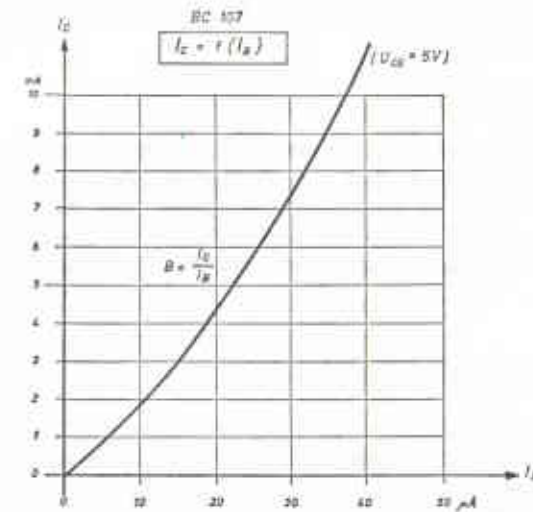


Abb. 4.22 — Strom-Steuerkennlinie $I_C = f(I_B)$ für einen Transistor BC 107

Dazu sei noch folgendes erwähnt: Der grundsätzliche Unterschied zwischen einer Spannungssteuerung und einer Stromsteuerung ergibt sich aus der Art der Steuerstromquelle. Besitzt die Steuerstromquelle einen vernachlässigbar **kleinen Innenwiderstand**, so ist ihre Ursprungsspannung ohne merklichen Verlust an den Klemmen und somit auch am Eingang Basis-Emitter des Transistors voll wirksam. Hier handelt es sich um **Spannungssteuerung**. Ist der Innenwiderstand der Steuerstromquelle dagegen **groß**, so tritt an ihm ein stromabhängiger Spannungsverlust $I \cdot R_i$ auf, um den die Klemmenspannung kleiner wird. Steigt die Ursprungsspannung, so steigt zwar der Strom, aber die Klemmenspannung weitaus weniger. Am Transistoreingang sind daher im wesentlichen nur die Stromänderungen wirksam; der Transistor wird **stromgesteuert**. Eine Stromsteuerung läßt sich durch künstliches Vergrößern des Stromkreiswiderstands erreichen, indem man zwischen Steuerstromquelle und Transistoreingang — also **in die Basiszuleitung** — einen **Relhenwiderstand schaltet**. Von dieser Möglichkeit wird sehr häufig Gebrauch gemacht.

Aus der Steigung der Steuerkennlinie läßt sich der Gleichstromverstärkungsfaktor des Transistors ermitteln. Je steiler die Steuerkennlinie verläuft, desto größer ist der Verstärkungsfaktor. Ebenso ist für jeden beliebigen Arbeitspunkt der zugehörige Gleichstromverstärkungsfaktor zu ermitteln, indem man den Kollektorstrom durch den

zugehörigen Basisstrom teilt ($\beta = \frac{I_C}{I_B}$).

Die Strom-Steuerkennlinie ermöglicht auch die Bestimmung des Wechselstromverstärkungsfaktors β . Dazu wird für eine festgelegte Basisstromänderung ΔI_B aus der Kennlinie die sich ergebende Kolle-

torstromänderung ermittelt. Aus dem Verhältnis ergibt sich der Wechselstromverstärkungsfaktor:

$$\beta = \frac{\Delta I_C}{\Delta I_B}$$

Die Steuerkennlinien eines Transistors kennzeichnen seine Verstärkerwirkung. Sie zeigen, daß eine Spannungssteuerung zu größeren Verzerrungen führt als eine Stromsteuerung. Aus der Strom-Steuerkennlinie lassen sich für jeden beliebigen Arbeitspunkt der Gleichstromverstärkungsfaktor und der Wechselstromverstärkungsfaktor ermitteln.

4.4. Der Transistor als Verstärker

In den bisherigen Schaltungen haben wir nur den Transistor selbst berücksichtigt und ihn — wie es in der Fachsprache heißt — im Kurzschluß betrieben. Dabei wurden immer konstante Gleichspannungen angelegt, so daß auch konstante Gleichströme über den Transistor flossen. Wie aber können wir die Eigenschaften des Transistors praktisch ausnutzen?

Nehmen wir einmal an, die Verstärkereigenschaft des Transistors soll dazu benutzt werden, schwache Tonfrequenzwechselströme zu verstärken. Die schwachen Tonfrequenzwechselströme könnten von einem dynamischen Mikrofon, einem einfachen Dioden-Rundfunkempfänger oder einem Tonabnehmer erzeugt werden. Ihre Stärke ist meistens so gering, daß sie nicht einmal in einem empfindlichen Kopfhörer hörbar zu machen sind. Wenn der Transistor schwache Wechselströme verstärken soll, so müssen wir ihn damit ansteuern und um die Verstärkerwirkung auszunutzen, einen Verbraucher einschalten. Unsere Schaltung sieht dann so aus:

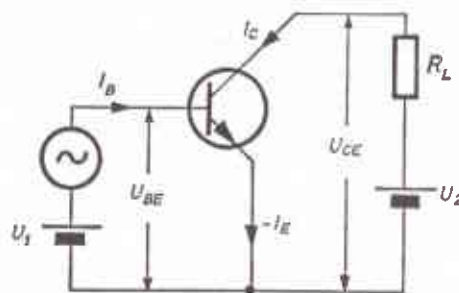


Abb. 4.23 — Der Transistor als Wechselstromverstärker

Im Basis-Stromkreis liegt eine Gleichstromquelle U_1 in Reihe mit dem Wechselstromgenerator (Mikrofon, Tonabnehmer). Ebenso benötigen wir zum Betrieb des Transistors eine Gleichstromquelle U_2 im Kollektorstromkreis. Beide Gleichstromquellen legen für den Transistor den Arbeitspunkt fest. In den Kollektorstromkreis wird der Verbraucher — im

einfachsten Fall ein ohmscher Widerstand — geschaltet. Der Verbraucherwiderstand wird allgemein als Lastwiderstand R_L bezeichnet. Der Verbraucher kann auch ein Kopfhörer, ein Lautsprecher oder eine weitere Verstärkerstufe sein. Infolge der angelegten Gleichspannungen fließt ein ganz bestimmter Basisruhestrom I_B und ebenso ein ganz bestimmter Kollektorruhestrom I_C . Liefert die im Basis-Emitter-Stromkreis liegende Wechselstromquelle jetzt einen Wechselstrom, so ändert sich der Basisstrom im Rhythmus dieses Wechselstroms; d.h., der Basisruhestrom wird mit Wechselstrom überlagert.

Bei sinusförmigem Wechselstrom zeigt das entsprechende Strom-Diagramm folgenden Verlauf:

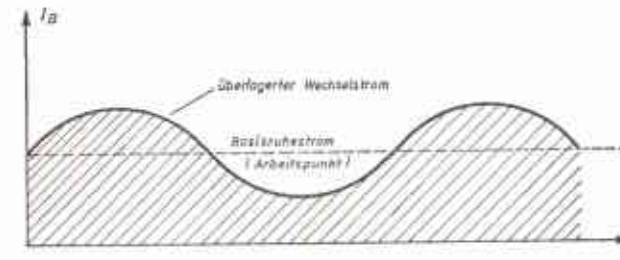


Abb. 4.24 — Zeitlicher Verlauf des wechselstromüberlagerten Basisstroms

Durch die Änderungen des Basisstroms wird aber — wie wir bereits festgestellt haben — eine viel stärkere (z.B. 100fache) Änderung des Kollektorstroms hervorgerufen. Das Strom-Zeit-Diagramm für den Kollektorstrom sieht also folgendermaßen aus:

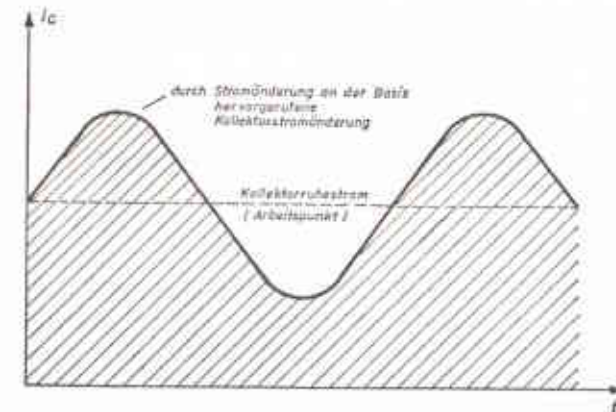


Abb. 4.25 — Zeitlicher Verlauf des Kollektorstroms als Folge der Basisstromänderungen (nicht maßstäblich!)

Beide Strom-Zeit-Diagramme lassen auch deutlich erkennen, wie wichtig die richtige Wahl des Arbeitspunkts ist. Bei zu niedrig gewähltem Basisruhestrom würden die negativen Halbwellen des überlagerten Wechselstroms unter Umständen in den negativen Bereich des Diagramms kommen. Eine Umkehrung der Basisstrom-Richtung hat aber eine Sperrung des Transistors zur Folge (PN-Übergang B—E). Dagegen bewirkt ein zu hoher Basisruhestrom zusammen mit den positiven Halbwellen des überlagerten Wechselstroms eine vielleicht zu hohe Stromspitze und damit einen zu starken Kollektorstrom. Die zulässige Verlustleistung könnte für den Transistor überschritten werden.

Diese Erkenntnis ist ebenfalls wichtig für den Bau von mehrstufigen Verstärkern. Da ein Transistor bei niedrigen Kollektorströmen weniger Rauschen erzeugt und am Eingang eines Verstärkers nur kleine Wechselspannungen oder Wechselströme liegen, wird für Anfangsstufen ein Kleinleistungstransistor mit niedriger Arbeitspunkteinstellung gewählt. Da die Wechselstromamplitude von Verstärkerstufe zu Verstärkerstufe größer wird, muß der Arbeitspunkt für die folgenden Stufen jeweils höher gewählt werden. Zwangsläufig damit müssen die Transistoren eine höhere zulässige Verlustleistung aufweisen als Anfangsstufen-Transistoren.

Kehren wir nach diesen Zwischenbetrachtungen zu unserer Verstärkerschaltung zurück: Die durch die Wechselstromquelle im Basisstromkreis bewirkten geringen Basisstromänderungen haben große Kollektorstromänderungen zur Folge. An dem im Kollektorstromkreis liegenden Verbraucher R_L werden die größeren Kollektorstromänderungen wirksam.

Die Kollektorstromänderungen bewirken eine entsprechende Spannungsänderung am Lastwiderstand:

$$\Delta U_L = \Delta I_C \cdot R_L.$$

4.4.1. Die wichtigsten Eigenschaften eines Transistors als Verstärker

4.4.1.1. Eingangswiderstand und Ausgangswiderstand

Wir haben bereits festgestellt, daß man zwischen dem statischen und dem dynamischen Widerstand eines Halbleiter-Bauelements unterscheiden muß. Das hängt mit der grundsätzlich gekrümmten Kennlinie eines Halbleiter-Bauelements zusammen. Für einen ohmschen Widerstand mit linearer Kennlinie sind dagegen statischer und dynamischer Widerstand gleich groß. Der Innenwiderstand einer Verstärkerstufe läßt sich wie jeder andere Widerstand anhand einer Span-

nungs- und Strommessung im Betriebszustand bestimmen. Man unterscheidet zwischen dem Gleichstromwiderstand R und dem Wechselstromwiderstand Z (auch als Impedanz bezeichnet).

Bestimmung R am Eingang: Dazu müssen die Basis-Emitter-Spannung und der Basisruhestrom bekannt sein. Aus beiden Werten ergibt sich

$$R_{\text{ein}} = \frac{U_{\text{BE}}}{I_B}.$$

Die Werte können z.B. für einen festgelegten Arbeitspunkt der Eingangskennlinie des Transistors entnommen werden.

Bestimmung R am Ausgang: Gemessen wird die am Transistor liegende Gleichspannung U_{CE} und der Kollektorstrom I_C .

$$\text{Dann ist } R_{\text{aus}} = \frac{U_{\text{CE}}}{I_C}.$$

Da der Transistor nahezu ausschließlich für Wechselstromvorgänge verwendet wird, interessiert bei ihm im wesentlichen auch der dynamische (Wechselstrom-)Widerstand Z am Eingang und Ausgang einer Transistorschaltung. Er läßt sich bestimmen zu

$$\begin{aligned} \text{Eingangswiderstand (dynamisch)} &= \frac{\text{Eingangswechselspannung}}{\text{Eingangswechselstrom}}, \\ \text{Ausgangswiderstand (dynamisch)} &= \frac{\text{Ausgangswechselspannung}}{\text{Ausgangswechselstrom}}. \end{aligned}$$

wobei die Spannungen direkt zwischen den Transistoranschlüssen zu messen sind.

Bestimmung Z am Ausgang: Der Transistor wird mit sinusförmigem Wechselstrom 1000 Hz angesteuert. Gemessen wird der am Transistor auftretende Wechselspannungsanteil von U_{CE} und der Wechselstromanteil von I_C .

$$\text{Dann ist } Z_{\text{aus}} = \frac{U_{\text{CE}} \sim}{I_C \sim}.$$

4.4.1.2. Die Stromverstärkung des Transistors

Je nach Anwendungszweck des Transistors als Verstärker kann man eine Verstärkerschaltung für optimale Stromverstärkung, Spannungsverstärkung oder Leistungsverstärkung auslegen. Zur Bestimmung der Stromverstärkung setzt man für die Emitterschaltung die erzielte Kollektorstromänderung der Basisstromänderung ins Verhältnis.

$$\beta = \frac{\Delta I_C}{\Delta I_B}.$$

Ebenso ist bekannt, daß der Wechselstromverstärkungsfaktor β nahezu gleich dem Gleichstromverstärkungsfaktor B eines Transistors ist.

In einem geschlossenen Stromkreis ist die Stromstärke um so größer, je kleiner der Stromkreiswiderstand und je höher die Spannung ist. Bei gegebener Versorgungsspannung ist demnach eine Stromerhöhung durch Verringern des Widerstands möglich. Daraus folgt: **Je größer die Stromverstärkung einer Transistor-Verstärkerstufe bei gegebener Versorgungsspannung werden soll, desto kleiner muß der Verbraucherwiderstand R_L gemacht werden.**

4.4.1.3. Die Spannungsverstärkung des Transistors

Anders als für große Stromverstärkung sind die Bedingungen für große Spannungsverstärkung. Ändern wir die Basis-Emitter-Spannung eines Transistors beispielsweise durch Überlagern der festen Spannung U_{BE} mit sinusförmiger Wechsellspannung, so treten im geschlossenen Basisstromkreis Stromschwankungen auf. Infolge der Verstärkerwirkung ergeben sich im Kollektorstromkreis stärkere Stromschwankungen. Da im Kollektorstromkreis ein Verbraucherwiderstand R_L liegt, treten an ihm entsprechende Spannungsschwankungen auf:

$$\Delta U_{RL} = \Delta I_C \cdot R_L.$$

Will man jetzt größere Spannungsänderungen ΔU_{RL} erzielen, so ist das nur durch Vergrößern des Verbraucherwiderstands R_L möglich; denn die Stromänderung ΔI_C liegt durch den Verstärkungsfaktor fest. Dabei dürfen wir jedoch eines unter keinen Umständen unberücksichtigt lassen: Je größer der Widerstand R_L ist, desto größer wird bei einem bestimmten Kollektorstrom I_C der Spannungsabfall an R_L . Um diesen Spannungsabfall ist aber die Spannung am Transistor (U_{CE}) kleiner. Sie kann so klein werden, daß der Transistor nicht mehr richtig arbeitet (wenn nämlich U_{CE} kleiner als $U_{CE\text{Rest}}$ wird). Der Spannungsverlust an R_L kann notfalls durch Erhöhen der Versorgungsspannung (Batteriespannung) aufgewogen werden.

Zusammenfassend ist hierzu also festzustellen: **Die Spannungsverstärkung ist für eine Transistor-Verstärkerstufe um so größer, je größer der Verbraucherwiderstand R_L im Kollektorstromkreis gemacht werden kann. Durch einen großen Verbraucherwiderstand wird ein Spannungsabfall hervorgerufen, um den die Spannung am Transistor kleiner wird. Größere Spannungsverstärkung bedingt daher eine höhere Versorgungsspannung (Batteriespannung).**

4.4.1.4. Die Leistungsverstärkung des Transistors

Eine optimale Leistungsverstärkung ist im allgemeinen für die letzte Stufe eines Transistor-Verstärkers wünschenswert, weil sie z.B. einen Lautsprecher mit ausreichender Wechselstromleistung versorgen muß. Die Leistung ergibt sich aus dem Produkt Spannung mal Strom ($P = U \cdot I$). Sie ist also um so größer, je größer Spannung **und** Strom gemacht werden können. Das bedeutet für einen Transistor-Verstärker sowohl optimale Spannungsverstärkung als auch optimale Stromverstärkung. Die Bedingungen dafür stehen sich aber entgegen. Eine große Leistung läßt sich daher bei gegebenem Verbraucherwiderstand nur mit größerer Versorgungsspannung erzielen. Dabei sind aber unter allen Umständen die zulässige Höchstspannung U_{CE} und die zulässige Verlustleistung zu beachten. Am einfachsten erzielt man größte Leistung durch **Anpassung**.

Zum besseren Verständnis eine etwas ausführlichere Erklärung des Begriffs Anpassung. Wählen wir der Einfachheit halber eine Batterie als Stromquelle, sie liefert eine bestimmte Leerlaufspannung U_0 (auch Ursprungsspannung genannt), besitzt aber auch einen inneren Widerstand R_i . Schalten wir jetzt einen Verbraucherwiderstand R_a an die Batterie, so tritt infolge des fließenden Stroms auch am inneren Widerstand der Batterie ein Spannungsabfall U_{vi} auf. An den Klemmen der Batterie — also auch am Verbraucherwiderstand R_a — liegt somit eine um den inneren Spannungsverlust geringere Spannung, die sogenannte Klemmenspannung: $U_K = U_0 - U_{vi}$.

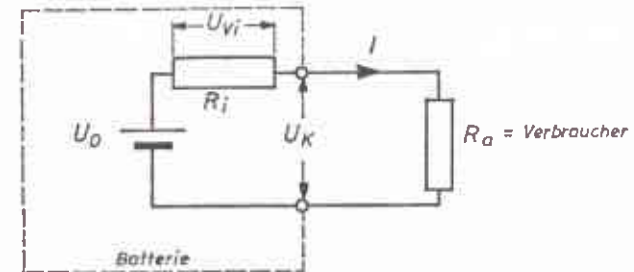


Abb. 4.26

Am anschaulichsten liegen die Verhältnisse leistungsmäßig bei den beiden Grenzfällen.

Grenzfall 1: Verbraucherwiderstand $R_a = 0 \text{ Ohm}$

In diesem Fall fließt der größtmögliche Strom (vgl. Abschn. 4.4.1.2). Wenn aber der Widerstand den Wert 0 Ohm hat, dann beträgt die Spannung an ihm 0 Volt (Kurzschluß!). Die Leistung an R_a beträgt $P_a = 0 \text{ Volt} \cdot I \text{ Ampere}$, also 0 Watt.

Grenzfall 2: Verbraucherwiderstand $R_a = \infty \text{ Ohm}$

In einem Stromkreis mit unendlich großem Widerstand fließt kein Strom. Da somit auch am inneren Widerstand der Batterie kein Spannungsabfall auftreten kann, liegt an den Batterieklemmen und am Verbraucherwiderstand jetzt die volle Leerlaufspannung U_0 . Die Leistung an R_a beträgt $P = U_0 \text{ Volt} \cdot 0 \text{ Ampere}$, also ebenfalls 0 Watt.

Für einen bestimmten Verbraucherwiderstand, dessen Wert zwischen 0 Ohm und ∞ Ohm liegt, ergibt sich eine größtmögliche Leistung. Das ist der Fall, wenn der Verbraucherwiderstand R_a genau gleich dem inneren Widerstand R_i der Batterie ist. Diesen Belastungsfall bezeichnet man als Anpassung.

Würden wir an die Batterie nacheinander Verbraucherwiderstände zwischen 0 Ohm und ∞ Ohm schalten und für jeden Verbraucherwiderstand die Leistung bestimmen, so ergibt sich anhand der Ergebnisse folgendes Diagramm:

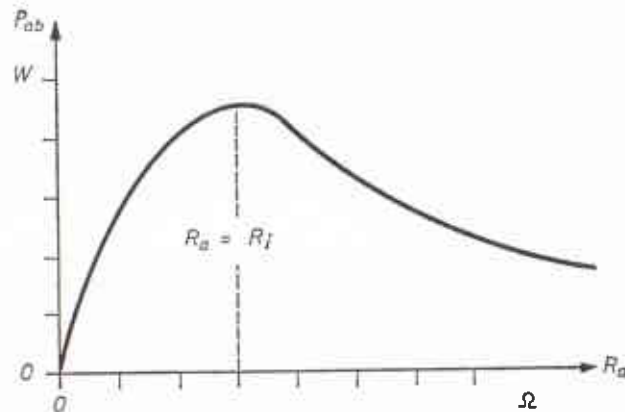


Abb. 4.27 — Anpassungsdiagramm

Fassen wir zusammen: Hohe Leistungsverstärkung setzt sowohl hohe Spannungs- als auch hohe Stromverstärkung voraus. Die größtmögliche Leistungsverstärkung ergibt sich, wenn der angeschlossene Verbraucherwiderstand R_a den gleichen Widerstand wie der Innenwiderstand R_i des Verstärkers hat. Für Wechselstromvorgänge gilt sinngemäß der entsprechende Wechselstromwiderstand Z_a bzw. Z_i .

Das Anpassungsdiagramm zeigt auch, daß sich eine Überanpassung ($R_a > R_i$) nicht so ungünstig auf die abgegebene Leistung auswirkt wie eine Unteranpassung ($R_a < R_i$). Deshalb soll man im Zweifelsfall R_a lieber etwas größer wählen; z.B. kann an einen Transistor-Leistungsverstärker wohl ein Lautsprecher mit größerem Widerstand als vorgesehen, dagegen nicht mit kleinerem angeschlossen werden.

Um Verzerrungen zu vermeiden, wird bei Leistungsverstärkern meistens R_i kleiner gewählt als der Verbraucherwiderstand.

4.4.1.5. Frequenzverhalten/Grenzfrequenz

Das Frequenzverhalten eines Transistors wird hauptsächlich durch die Kapazität zwischen Basis und Emitter und dem parallelliegenden

ohmschen Widerstand am Eingang bestimmt. Zur Bestimmung der Kapazität zwischen Basis und Emitter kann man für Niederfrequenz-Leistungstransistoren folgende Erfahrungsformel überschläglich anwenden:

$$C_{BE} = 8000 \cdot I_C \text{ pF, wobei } I_C \text{ in mA}$$

einzusetzen ist. Bei einem Kollektorstrom von 1 mA beträgt also die Kapazität zwischen Basis und Emitter rd. 8000 pF. Diese Kapazität stellt für Wechselströme einen Nebenschluß dar, der sich mit steigender Frequenz zunehmend nachteilig auswirkt. Liegt jetzt parallel zu dieser Kapazität ein ohmscher Widerstand, so hängt der Einfluß der Kapazität von der Größe des ohmschen Widerstands, genauer vom Größenverhältnis des ohmschen Widerstands zum kapazitiven Blindwiderstand ab.

Beispiel:

Der Arbeitspunkt eines NF-Transistors liegt bei $I_C = 2$ mA. Sein Eingangswiderstand beträgt 1000 Ω . Für $I_C = 2$ mA ergibt sich eine Basis-Emitter-Kapazität von $C_{BE} = 8000 \cdot 2 = 16000$ pF = 16 nF. Die Schaltung sieht dann so aus:

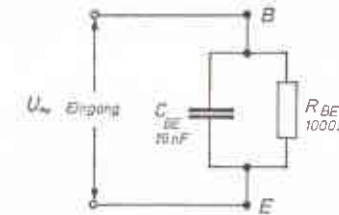


Abb. 4.28

Legen wir an die Schaltung eine Wechselspannung steigender Frequenz, so ergeben sich bei 1000 Hz folgende Verhältnisse:

$$X_C = \frac{1}{\omega \cdot C_{BE}} = \frac{1}{2 \pi \cdot 1000 \cdot 16 \cdot 10^{-9}} = \frac{10^9}{2 \pi \cdot 1000 \cdot 16} \approx 10\,000 \, \Omega.$$

Parallel zum 1000-Ohm-Widerstand liegt also der Blindwiderstand 10 000 Ohm des Kondensators C_{BE} . Da er 10mal so groß ist, wirkt er sich kaum aus.

Anders liegen aber bereits die Verhältnisse bei 10 000 Hz. Dafür ergibt sich

$$X_C = \frac{1}{\omega \cdot C_{BE}} = \frac{1}{2 \pi \cdot 10\,000 \cdot 16 \cdot 10^{-9}} = \frac{10^9}{2 \pi \cdot 10\,000 \cdot 16} \approx 1000 \, \Omega.$$

Bei 10 000 Hz ist also der kapazitive Blindwiderstand gerade noch so groß wie der ohmsche Parallelwiderstand. Der Kondensator C_{BE} stellt schon einen erheblichen Nebenschluß dar. Würde der ohmsche Widerstand aber nur einen Wert von 10 Ω besitzen, dann hätte die Kapazität C_{BE} erst bei 1 000 000 Hz = 1 MHz den Blindwiderstand von 10 Ω . Ähnliche Verhältnisse ergeben sich bei der Basisschaltung.

An diesen Zahlenbeispielen erkennen wir den großen Einfluß des Eingangswiderstands und der Kapazität C_{pg} auf das Frequenzverhalten des Transistors.

Auch zwischen den anderen Schichten des Transistors sind Kapazitäten wirksam (also zwischen C und B bzw. zwischen C und E); sie sind aber für Niederfrequenzen (Tonfrequenzen) vernachlässigbar klein.

Die Frequenz, bei der $X_{CBE} = R_{BE}$ ist, wird als **Grenzfrequenz** bezeichnet. Unter dieser Bedingung liegt am Eingang noch der 0,7fache Wert des Steuerstroms. Die Grenzfrequenz wird für Transistoren in den Datenblättern angegeben und je nach Schaltungsart, für die sie gilt, wie folgt bezeichnet:

- Grenzfrequenz für Basisschaltung f_a ,
- Grenzfrequenz für Emitterschaltung f_β ,
- Grenzfrequenz für Kollektorschaltung f_y (wenig gebräuchlich).

Unter der Grenzfrequenz versteht man für eine bestimmte Transistor-schaltung diejenige Frequenz, bei der der Wechselstromverstärkungs-faktor auf den 0,7fachen Wert gegenüber der Verstärkung bei niedri-gen Frequenzen abgesunken ist.

4.4.2. Die drei Grundschaltungen für einen Transistor

Anhand unserer bisherigen Betrachtungen über die Schaltung eines Transistors mußte man den Eindruck gewinnen, daß es nur eine ein-zige grundsätzliche Schaltungsart gibt. Bei den behandelten Transi-stor-Schaltungen war immer der Emitteranschluß gemeinsame Elek-trode für den Eingang (Steuerstromkreis) und den Ausgang (Last-oder Verbraucherstromkreis) eines Transistorverstärkers. Es besteht jedoch grundsätzlich die Möglichkeit, auch den Basis- oder Kollektor-an-schluß des Transistors als gemeinsame Elektrode für den Ein- und Ausgang eines Verstärkers zu schalten.

Bei diesen Schaltungen ist zu unterscheiden zwischen:

- der **Emitterschaltung**, bei der der Emitteranschluß die ge-meinsame Elektrode für Ein- und Ausgang darstellt,
- der **Basisschaltung**, bei der der Basisanschluß als gemein-same Elektrode für Ein- und Ausgang gilt, und
- der **Kollektorschaltung**, bei der sinngemäß der Kollektor-an-schluß die gemeinsame Elektrode für Ein- und Ausgang ist.

Die drei Grundschaltungen sind in Abb. 4.29 und 4.30 dargestellt. Dazu ist noch zu bemerken, daß die Batterie oder auch das Netzgerät für den Wechselstrom grundsätzlich einen Kurzschluß darstellt. Er ist durch eine gestrichelte Kapazität in den Schaltbildern angedeutet.

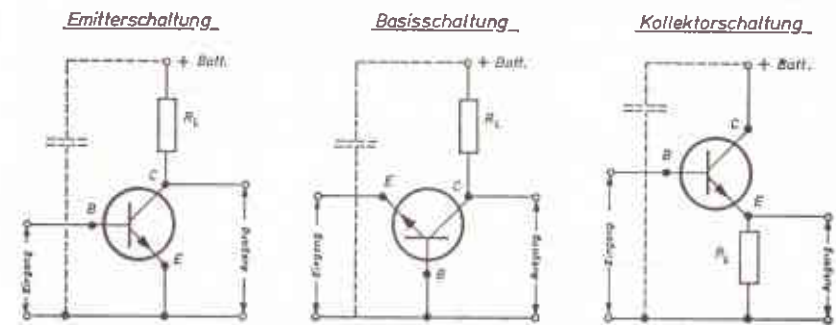


Abb. 4.29 — Die Grundschaltungen für einen NPN-Transistor

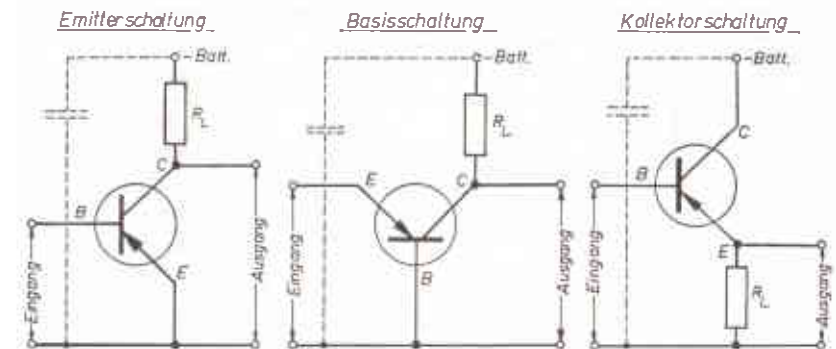


Abb. 4.30 — Die Grundschaltungen für einen PNP-Transistor

Wie aus den beiden Abbildungen ersichtlich, unterscheiden sich NPN-und PNP-Transistoren nur durch die Polarität (Batteriepolung). Die Eigenschaften der drei Grundschaltungen können der Tabelle auf Seite 154 entnommen werden.

Die in der Tabelle angegebenen Werte sind als Richtwerte — also größenordnungsmäßig — anzusehen. Aus der Gegenüberstellung der drei Schaltungsarten geht eindeutig hervor, daß die Emitterschaltung für allgemeine Verstärker am günstigsten ist, da sie sowohl eine große Stromverstärkung als auch eine große Spannungsverstärkung und somit die Voraussetzungen für eine hohe Leistungsverstärkung ergibt. Müssen Wechselspannungen sehr hoher Frequenz verstärkt werden, so wird die Basisschaltung angewandt.

	Emitterschaltung	Basisschaltung	Kollektorschaltung
Eingangswiderstand Z_{ein}	ca. 100 Ω mittel bis 10 000 Ω	ca. 10 Ω klein bis 100 Ω	ca. 10 000 Ω groß bis 100 000 Ω
Ausgangswiderstand Z_{aus}	ca. 1 000 Ω mittel bis 10 000 Ω	ca. 10 000 Ω groß bis 100 000 Ω	ca. 10 Ω klein bis 100 Ω
Stromverstärkung α β γ	groß $\beta = \frac{\Delta I_C}{\Delta I_B}$ (10- bis 500fach)	< 1 $\alpha = \frac{\Delta I_C}{\Delta I_E}$ (kleiner als 1)	groß $\gamma = \frac{\Delta I_E}{\Delta I_B}$ (10- bis 500fach)
Spannungsverstärkung v_u	groß $v_{ue} = \frac{\Delta U_{CE}}{\Delta U_{BE}}$ (100- bis 1000fach)	groß $v_{ub} = \frac{\Delta U_{CB}}{\Delta U_{BE}}$ (100- bis 1000fach)	< 1 $v_{uc} = \frac{\Delta U_{CE}}{\Delta U_{CB}}$ (etwas kleiner als 1)
Leistungsverstärkung v_p	sehr groß $v_{pe} = \beta \cdot v_{ue}$	mittel $v_{pb} = \alpha \cdot v_{ub}$	klein $v_{pc} = \gamma \cdot v_{uc}$
Grenzfrequenz	mittel f_β bis ca. 10 MHz	hoch f_α bis ca. 5 GHz $f_\alpha \approx \beta \cdot f_\beta$	mittel f_γ bis ca. 10 MHz $f_\gamma \approx f_\beta$
Anwendung für	NF-Verstärker und HF-Verstärker, Leistungsverstärker, Schalter	HF-Verstärker für hohe Frequenzen	Anpassungsstufen (sog. Impedanzwandler)

In Basisschaltung liegt parallel zur Eingangskapazität C_{BE} nur ein kleiner Eingangswiderstand von ca. 10 bis 100 Ω . Somit wirkt sich der Blindwiderstand der Basis-Emitter-Kapazität des Transistors auch erst bei sehr viel höheren Frequenzen als Nebenschluß aus. (Vgl. Abschn. 4.4.1.5 „Frequenzverhalten/Grenzfrequenz“.) Man findet daher Transistorverstärkerstufen in Basisschaltung vor allem in den HF-Eingangsstufen von UKW-Rundfunk- und Fernsehgeräten.

Für die Kollektorschaltung ergibt sich ein sehr hoher Eingangs- und ein sehr kleiner Ausgangswiderstand. Man wendet sie daher zur Anpassung einer hochohmigen an eine niederohmige Stufe bzw. an einen niederohmigen Verbraucher an. Sie bietet gegenüber einem Anpassungsübertrager Vorteile hinsichtlich des geringeren Raumbedarfs und der geringeren Verzerrungen.

Die häufigste und vorteilhafteste Schaltung für allgemeine Verstärkeranwendungen ist die Emitterschaltung, da sie sowohl hohe Strom- als

auch Spannungsverstärkung und damit auch hohe Leistungsverstärkung ergibt.

Müssen Wechselspannungen sehr hoher Frequenz verstärkt werden, so ist die Basisschaltung anzuwenden. Für die Basisschaltung ergibt sich eine hohe Grenzfrequenz.

Die Kollektorschaltung wird zur Anpassung (Impedanzwandler) verwendet, da sie einen hohen Eingangs- und einen niedrigen Ausgangswiderstand besitzt.

Wir werden uns nun künftig ausschließlich mit der Emitterschaltung befassen, da sie die breiteste Anwendung findet. Sämtliche Schaltungen und Berechnungen beziehen sich auf NPN-Transistoren. Bei Verwendung von PNP-Transistoren sind die anzulegenden Spannungen umzupolen.

4.4.3. Die Stromversorgung von Transistorverstärkern

Bei der Stromversorgung von Transistorverstärkern unterscheiden wir im wesentlichen zwischen dem **Batteriebetrieb** und dem **Netzbetrieb**. Die Hauptmerkmale sowie Vor- und Nachteile dieser beiden Betriebsarten sollen hier kurz gegenübergestellt werden.

4.4.3.1. Batteriebetrieb

Vorteile: Geringer Schaltungsaufwand; verhältnismäßig geringes Gewicht. Unabhängig vom Netzanschluß, daher ideal für tragbare Geräte; somit auch überall betriebsbereit. Wegen des geringen Spannungsbedarfs für Transistorverstärker einfacher Batterieaufwand (geringe Zellenzahl).

Nachteile: Die Batteriespannung ist nicht konstant; sie sinkt bis zur zulässigen Entladegrenze auf den halben Wert der Nennspannung. Das bedingt u.U. zusätzlichen Aufwand für Stabilisierungsschaltungen. Erheblich höhere Betriebskosten gegenüber Netzbetrieb (1 kWh kostet bei Netzbetrieb etwa 11 Pf, bei Batteriebetrieb dagegen etwa zwischen 150 DM und 600 DM!).

Der Energiebedarf — vor allem bei kleinen Transistorgeräten — ist wegen des hohen Wirkungsgrads der Transistoren verhältnismäßig gering. Aus diesem Grunde arbeiten z.B. batteriebetriebene Transistorgeräte auch erheblich wirtschaftlicher als batteriebetriebene Röhrengeräte.

4.4.3.2. Netzbetrieb

Vorteile: Vor allem sehr wirtschaftlich. Durch entsprechende Schaltungsmaßnahmen sehr konstante Spannung möglich.

Nachteile: Wegen der Abhängigkeit vom Netz eben nur dort betriebsbereit, wo ein Netzanschluß vorhanden ist. Anschaffungskosten erheblich höher als für Batterien. Höheres Gewicht — schon allein wegen des erforderlichen Netztransformators. Höherer Schaltungsaufwand, da gleichgerichtet, gesiebt und stabilisiert werden muß.

Wir möchten in diesem Zusammenhang dringend vor der Herstellung einfacher Netzgeräte im „Eigenbau“ warnen, wenn nicht komplett mit Anschlußkabel und Stecker montierte und den **VDE-Vorschriften** entsprechende Transformatoren verwendet werden. Ebenso ist von der Anschaffung sehr billiger Netzgeräte dringend abzuraten, da diese häufig nicht den VDE-Vorschriften entsprechen und **keinen galvanischen Schutz gegenüber der Netzspannung bieten**.

4.4.3.3. Die Betriebsspannungen des Transistors

Aus den aufgenommenen Kennlinien war zu ersehen, daß zum Betrieb eines Transistors eine kleine Spannung von ca. 0,2 bis 0,7 Volt zwischen Basis und Emittor sowie eine Spannung von einigen Volt zwischen Kollektor und Emittor notwendig ist. Eine Erhöhung der Kollektor-Emittor-Spannung U_{CE} über $U_{CE\text{Rest}}$ ergab nur noch eine sehr geringe Steigerung des Kollektorstroms. Diese Tatsache führt sehr oft zu der irrtümlichen Annahme, zum Betrieb des Transistors genüge eine kleine Spannung von 1 bis 2 V. Jedoch darf dabei eines nicht außer acht gelassen werden:

Bei den bisherigen Meßschaltungen haben wir den Transistor grundsätzlich im Kurzschluß betrieben. Sobald aber die Verstärkereigenschaft eines Transistors ausgenutzt werden soll, muß an den Eingang eine Steuerstromquelle und in den Ausgangstromkreis ein Verbraucher (Lastwiderstand) geschaltet werden. Vor allem der Lastwiderstand im Kollektorstromkreis verursacht einen Spannungsabfall, um den die eigentliche Spannung U_{CE} geringer wird. Je größer der Lastwiderstand und der Kollektorstrom sind, desto größer wird der Spannungsabfall. Für eine hohe Spannungsverstärkung ist aber ein hoher Lastwiderstand wünschenswert.

Die Versorgungsspannung muß um den am Lastwiderstand auftretenden Spannungsabfall größer gemacht werden. Für Transistorverstärker sind daher Versorgungsspannungen (Batterie- oder Netzgerät) zwischen etwa 6 V und 50 V üblich. Wir verwenden der Einfachheit halber zwei in Reihe geschaltete 4,5-V-Flachbatterien und erhalten damit eine Spannung von 9 V.

Woher bekommen wir nun die Spannung zwischen Basis und Emittor? Dafür ist nicht etwa — wie man zunächst annehmen könnte — eine zweite Spannungsquelle notwendig. Wir können nämlich die Basis-Emittor-Spannung ebenfalls aus der 9-V-Batterie gewinnen; denn die Richtung der Kollektor-Emittor-Spannung und der Basis-Emittor-Spannung ist gleich. Jetzt kann der Vorteil genutzt werden, der sich aus der abgeänderten Schaltung am Schluß des Abschn. 4.1 ergibt.

Zur Erzeugung der Basis-Emittor-Spannung aus der Versorgungsspannung eines Transistorverstärkers gibt es zwei grundsätzliche Schaltungsmaßnahmen:

a) Schaltung mit Basisvorwiderstand

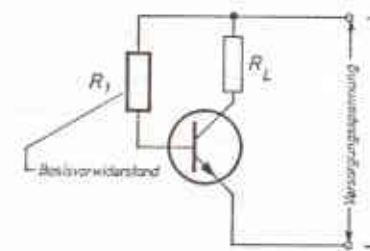


Abb. 4.31

Dieses sehr einfache Verfahren wird vor allem für Vorverstärker angewendet, wenn keine hohen Anforderungen an die Stabilität des Arbeitspunkts gestellt werden. Der Basisvorwiderstand stellt in Reihenschaltung mit dem Durchlaßwiderstand Basis-Emittor einen Spannungsteiler dar. Wir können ihn so bemessen, daß zwischen Basis und Emittor genau die benötigte Betriebsspannung U_{BE} liegt. Auf die Berechnung des Widerstandswerts für den Vorwiderstand kommen wir noch zurück.

b) Schaltung mit Basis-Spannungsteiler

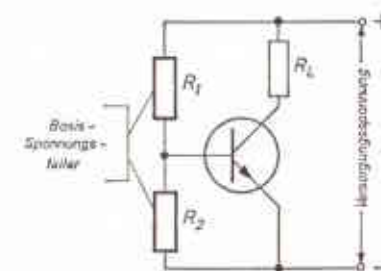


Abb. 4.32

Mit Hilfe dieses Spannungsteilers wird an R_2 eine Teilspannung erzeugt, die dem erforderlichen Wert für U_{BE} entspricht. Der Spannungsteiler wird so bemessen, daß er selbst einen ca. 5—10mal höheren Strom, als der Basisstrom aufnimmt. Dadurch wird eine verhältnismäßig gute Stabilisierung des Arbeitspunkts erreicht. Wir brauchen nur zu berücksichtigen, daß der Widerstand von Halbleitern mit stei-

gender Temperatur kleiner wird, um zu erkennen, daß dadurch eine Änderung der Teilspannung auftreten muß. Damit würde sich aber der Arbeitspunkt verschieben. Dieser Nachteil wird durch den 5- bis 10fachen Basisstromwert im Basis-Spannungsteiler weitgehend vermieden. Der Basis-Spannungsteiler ist für die Stabilisierung des Arbeitspunkts jedoch nur dann voll wirksam, wenn gleichzeitig Stromgegenkopplung (Widerstand in der Emitterzuleitung) angewandt wird.

Fassen wir zusammen: **Die zum Betrieb eines Transistors erforderlichen Spannungen U_{CE} und U_{BE} können an einer und derselben Spannungsquelle entnommen werden.** Dabei wird die kleinere Basis-Emitter-Spannung U_{BE} an einem Spannungsteiler abgegriffen. Berücksichtigt man beim Batteriebetrieb, daß die Batteriespannung im Laufe der Entladung auf den halben Wert der Nennspannung sinkt, so kann der Arbeitspunkt der Transistorschaltung für etwa $\frac{1}{2}$ der Batterie-Nennspannung ausgelegt werden. Dann liegt der Arbeitspunkt bei frischer Batterie etwas zu hoch und bei entladener Batterie etwas zu niedrig gegenüber der günstigsten Einstellung. Dabei ist wichtig, daß der Arbeitspunkt bei frischer Batterie noch unterhalb der Leistungshyperbel der für den Transistor zulässigen Verlustleistung liegt.

4.4.4. Die Ein- und Auskopplung des Signals

Für die Ein- und Auskopplung des Signals an einem Transistorverstärker unterscheidet man hauptsächlich drei Schaltungsarten (C-Kopplung, Transformatorkopplung, Gleichstromkopplung). Da die Steuerstromquellen, z.B. Mikrofon, Tonabnehmer, Rundfunkempfänger usw., einen bestimmten Innenwiderstand besitzen, verändern sie bei direktem Anschluß an den Eingang einer Transistorstufe die Arbeitspunkteinstellung.

Wir brauchen uns nur vorzustellen, daß der Widerstand R_2 eines Basis-Spannungsteilers den Wert 10 kOhm hat. Wird parallel dazu ein Mikrofon mit 200 Ohm Innenwiderstand geschaltet, so ist der Gesamtwiderstand kleiner als 200 Ohm. Damit sinkt aber die Teilspannung auf etwa $1/50$ des ursprünglichen Werts ab. Es liegt auf der Hand, daß der Transistor mit $1/50$ der erforderlichen Basis-Emitter-Spannung nicht mehr arbeitet.

Die Schaltungen zur Ein- und Auskopplung eines Wechselstromsignals sind daher so auszulegen, daß der vorgesehene Arbeitspunkt des Transistors erreicht bzw. gehalten wird.

4.4.4.1. Die C-Kopplung (auch RC-Kopplung)

Die einfachste und häufigste Schaltungsmaßnahme zur Ein- und Auskopplung eines Wechselstromsignals ist die C-Kopplung, genauer Kapazitätskopplung. Sie wird häufig auch als RC-Kopplung bezeichnet. Schalten wir in die Basiszuleitung vor den Basis-Spannungsteiler einen

Kondensator, so trennt dieser die Transistorschaltung gleichstrommäßig von der Steuerstromquelle. Es ist jetzt grundsätzlich einerlei, welchen Innenwiderstand die Steuerstromquelle selbst hat. Der durch den Spannungsteiler eingestellte Arbeitspunkt, d.h. die Basis-Emitter-Spannung, bleibt davon unbeeinflusst.

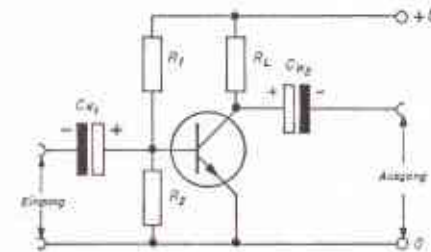


Abb. 4.33 — C-Kopplung bei einem Transistorverstärker

Das gleiche trifft auch für den Ausgang des Transistorverstärkers zu. Eine Änderung des Lastwiderstands durch direkte galvanische Reihen- oder Parallelschaltung eines Verbrauchers (z.B. Kopfhörer, Lautsprecher oder folgende Verstärkerstufe) würde den Arbeitspunkt verschieben. Auch hier bietet sich als einfachstes Schaltelement der Kondensator an. Er trennt den Ausgangsstromkreis gleichstrommäßig von einem angeschlossenen Verbraucher.

Soll mit dem Transistorverstärker ein breites Frequenzband übertragen werden, so ist der Kopplungskondensator nach der tiefsten zu übertragenden Frequenz zu bemessen; denn aufgrund der Beziehung

$X_C = \frac{1}{\omega \cdot C}$ wird der Wechselstromwiderstand des Kondensators um

so größer, je niedriger die Frequenz ist. Überschläglich kann der notwendige Kapazitätswert aus folgender Formel bestimmt werden:

$$R = \frac{1}{2\pi \cdot f_u \cdot C_K} \quad \text{bzw. umgestellt: } C_K = \frac{1}{2\pi \cdot f_u \cdot R}$$

In der Praxis wird der Wert für C_K etwa 3—5mal so groß gewählt wie der berechnete Wert.

Darin bedeuten:

R = der gesamte ohmsche Widerstand, der hinter der Kapazität liegt,

f_u = niedrigste zu übertragende Frequenz,

C_K = Kapazität des Kopplungskondensators in F.

Beispiel:

Mit einer Transistorstufe sollen Tonfrequenzen zwischen 30 Hz und 15 000 Hz verstärkt werden. Der Eingangswiderstand (Parallelschaltung aus R_2 des Basis-Spannungsteilers und dem Basis-Emitter-Widerstand R_{BE} des Transistors) beträgt 2 kOhm. Welchen Kapazitätswert muß der Kopplungskondensator haben?

Gegeben: $R = 2 \text{ k}\Omega$; $f = 30 \dots 15\,000 \text{ Hz}$, also $f_u = 30 \text{ Hz}$

Gesucht: C_K

$$\text{Lösung: } C_K = \frac{1}{2\pi \cdot f_u \cdot R} = \frac{1}{2 \cdot 3,14 \cdot 30 \cdot 2000} = \frac{1}{120\,000 \cdot 3,14} \text{ F}$$

$$C_K = 0,00000265 \text{ F} = 2,65 \text{ }\mu\text{F, gewählt } 8 \text{ }\mu\text{F.}$$

Der Kopplungskondensator müßte also eine Kapazität von rund $8 \text{ }\mu\text{F}$ haben, damit die tiefste Frequenz 30 Hz noch ausreichend übertragen wird. Nach der gleichen Formel kann übrigens auch die Kapazität C_E zur Überbrückung des Emitterwiderstands R_E berechnet werden (Abschn. 4.4.5.2). Da für die C-Kopplung im allgemeinen große Kapazitätswerte erforderlich sind, werden bevorzugt Elektrolytkondensatoren verwendet. **Beim Einsatz von Elektrolytkondensatoren ist unbedingt darauf zu achten, daß der Plusanschluß des Kondensators immer am höheren Pluspotential einer Schaltung liegt.** Im Zweifelsfall können die Potentiale in einer Schaltung durch eine einfache Messung bestimmt werden (vgl. hierzu Abschn. 1.2.1).

4.4.4.2. Die Transformatorkopplung (Übertragerkopplung)

Da ein Transformator zwei Stromkreise galvanisch trennt, bietet er sich ebenfalls als Koppelement an. Zudem hat man bei ihm noch den Vorteil, daß durch geeignete Wahl des Übersetzungsverhältnisses eine ideale Anpassung von Steuerstromquelle an den Transistoreingang und Verbraucher an den Transistorausgang erzielt wird. Mit Hilfe der Transformatorkopplung läßt sich daher der günstigste Wirkungsgrad eines Transistorverstärkers erreichen.

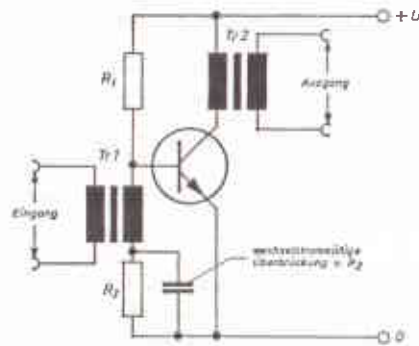


Abb. 4.34 — Transformatorkopplung bei einem Transistorverstärker

Nun hat ein Transformator aber leider neben seinem höheren Preis und seinem größeren Volumen auch noch einen weiteren Nachteil, den Eisenkern. Infolge der gekrümmten Magnetisierungskurve des Eisen-

kerns wird der Wechselstrom mit einem Transformator nicht verzerrungsfrei übertragen. Das spielt zwar in der Starkstromtechnik keine entscheidende Rolle; jedoch sind bei der Übertragung von Nachrichten oder Musik auch kleinste Verzerrungen unerwünscht. Daher geht man in der Transistorschaltungstechnik immer mehr zu anderen Kopplungsmaßnahmen über. Unter den technischen Daten eines Verstärkers oder Kofferradios findet man häufig den Hinweis „eisenlose Endstufe“. Er bedeutet, daß zur Ankopplung des Lautsprechers kein Transformator verwendet wird. Bei den für Kopplungszwecke hergestellten Transformatoren, die man eigentlich Übertrager nennt, bezieht sich das angegebene Übersetzungsverhältnis auf das Widerstandsübersetzungsverhältnis. So bedeutet z.B. die Angabe $1:10$, daß sich der Wechselstromwiderstand Z_1 der Primärwicklung zum Wechselstromwiderstand Z_2 der Sekundärwicklung wie $1:10$ verhält.

Da die Berechnung der Wickeldaten für einen Kopplungstransformator sehr umfangreich und schwierig ist und die Transformatorkopplung in der Transistor-Schaltungstechnik an Bedeutung verliert, wollen wir in diesem Rahmen auf Beispiele verzichten.

4.4.4.3. Die Gleichstromkopplung (galvanische Kopplung)

Diese Kopplungsart ergibt die optimalsten Eigenschaften für einen Transistorverstärker. Sie ist aber leider auch die schwierigste; denn hier müssen Innenwiderstand der Steuerstromquelle (z.B. Mikrofon, Tonabnehmer, Vorverstärker) sowie Innenwiderstand des Verbrauchers (z.B. Kopfhörer, Lautsprecher oder nachfolgende Verstärkerstufe) gleichstrommäßig in die Verstärkerschaltung einbezogen werden. An diesen Widerständen darf sich nichts ändern, weil sich sonst automatisch der Arbeitspunkt des Transistors verschiebt. Als Beispiel einer Gleichstromkopplung ist in Abb. 4.35 ein zweistufiger Verstärker dargestellt.

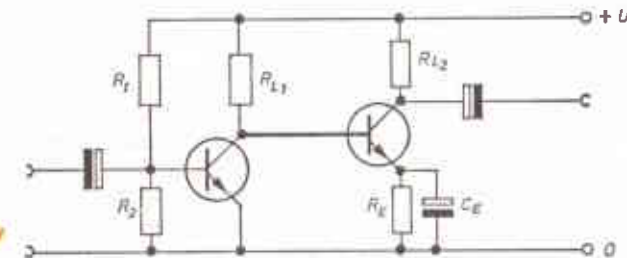


Abb. 4.35 — Gleichstromkopplung zwischen zwei Transistorverstärkerstufen

Es liegt auf der Hand, daß bei dieser Kopplungsart keine Wechselstromwiderstände (von Kapazitäten oder Induktivitäten) vorhanden sind. Somit kann durch die Gleichstromkopplung auch keine Frequenzbeeinflussung auftreten (vgl. hierzu auch die Berechnung des Kapazitätswerts für C-Kopplung). Im abgebildeten Schaltungsauszug müssen die Arbeitspunkte beider Transistoren so eingestellt werden, daß U_{CE} des ersten genau U_{BO} des zweiten Transistors ergibt. Liegt U_{CE} höher, so kann ein Widerstand in den Kopplungsweg geschaltet werden.

Die Gleichstromkopplung wird für Wechselstromverstärker mit möglichst linearem Frequenzgang, für Gleichstrom- oder Gleichspannungsverstärker und Kippstufen angewendet. Gleichstrom- oder Gleichspannungsverstärker haben große Bedeutung für die Meßtechnik, nämlich für den Bau von z.B. Transistorvoltmetern. Auch in dem später behandelten Verstärker wird zum Teil Gleichstromkopplung angewendet.

Zusammenfassend sollen hier noch einmal die wesentlichen Vor- und Nachteile der drei Kopplungsarten genannt werden:

a) C-Kopplung

Vorteile: Einfach zu verwirklichen. Die Transistorstufe kann gleichstrommäßig völlig unabhängig vom Innenwiderstand der Steuerstromquelle bzw. des Verbrauchers auf den gewünschten Arbeitspunkt eingestellt werden.

Nachteile: Eine Kapazität hat einen frequenzabhängigen Wechselstromwiderstand. Dadurch werden tiefe Frequenzen benachteiligt. Vor allem bei kleinem Eingangswiderstand und kleinem Verbraucherwiderstand werden große Kapazitätswerte benötigt, um die tiefste Frequenz noch befriedigend zu übertragen.

b) Trafokopplung

Vorteile: Durch Wahl eines geeigneten Übersetzungsverhältnisses ist eine ideale Anpassung möglich. Der Gleichstromwiderstand der Windungen ist klein. Dadurch kann große Strom- und Leistungsverstärkung erzielt werden.

Nachteile: Infolge der gekrümmten Magnetisierungskurve des Eisenkerns treten zusätzliche Verzerrungen auf. Der Transformator (Übertrager) ist teuer und beansprucht ein verhältnismäßig großes Bauvolumen; er ist außerdem schwer, was bei tragbaren Geräten sehr nachteilig ist. Jeweils eine Windung liegt galvanisch verbunden im Ein- bzw. Ausgangstromkreis. Zur Arbeitspunkteinstellung müssen daher die Gleichstromwiderstände dieser Windungen berücksichtigt werden.

c) Gleichstromkopplung

Vorteile: Im Kopplungsweg sind keine frequenzbeeinflussenden Bauelemente vorhanden. Es ist auch eine Verstärkung von reinen Gleichstromgrößen (Spannung, Strom) möglich, was z.B. bei Meßverstärkern gefordert wird. Keine zusätzlichen Bauelemente erforderlich.

Nachteile: Schwierig einzustellender Arbeitspunkt, da Steuerstromquelle bzw. Verbraucher galvanisch mit der Transistorstufe verbunden ist. Jede Änderung des Widerstands von Steuerstromquelle bzw. Verbraucher hat eine Verschiebung des Arbeitspunkts zur Folge.

Bei allen Kopplungsarten ist unbedingt darauf zu achten, daß der vorgesehene Arbeitspunkt für die Transistoren erreicht und im Betrieb nicht unbeabsichtigt verschoben werden kann.

4.4.5. Die Arbeitspunktstabilisierung der Transistorverstärker

Neben der Festlegung des günstigsten Arbeitspunkts für einen Transistor ist der Arbeitspunktstabilisierung besondere Aufmerksamkeit zu widmen. Bei der Besprechung der Halbleiter-Physik sowie durch Versuche haben wir festgestellt, daß das elektrische Leitvermögen eines Halbleiters bzw. eines Halbleiter-Bauelements sehr stark temperaturabhängig ist. So können bereits Temperaturschwankungen von $\pm 10^\circ\text{C}$ — wenn keine Stabilisierungsmaßnahmen angewendet werden — den Arbeitspunkt eines Transistors so weit verschieben, daß dieser nicht mehr einwandfrei arbeitet oder zerstört wird. **Da die Erwärmung von Transistoren vor allem durch die Verlustleistung verursacht wird und der Kollektorstrom die größte Verlustleistung bewirkt, dienen entsprechende Schaltungsmaßnahmen nahezu ausschließlich der Stabilisierung des Kollektorstroms.** Vergleichen wir einmal anhand einfacher Zahlenbeispiele:

Da beim Transistor $I_C \approx I_E$ ist, gilt für die durch den Kollektorstrom verursachte Verlustleistung:

$$P_{vC} = U_{CE} \cdot I_C = 5\text{ V} \cdot 10\text{ mA} = 50\text{ mW}.$$

Der Anteil der durch den Basisstrom bewirkten Verlustleistung ist dagegen so gering, daß er vernachlässigt werden kann. Da die Basisstromstärke nur etwa $\frac{1}{100}$ der Kollektorstromstärke beträgt, gilt:

$$P_{vB} = U_{BE} \cdot I_B = 0,6\text{ V} \cdot 0,1\text{ mA} = 0,06\text{ mW}.$$

Schaltungsmaßnahmen zur Stabilisierung des Kollektorstroms sind um so wichtiger, je näher der Arbeitspunkt eines Transistors an der zulässigen Grenze (Leistungshyperbel) liegt. Die Stabilisierungsschaltungen setzen jedoch den Wirkungsgrad eines Transistorverstärkers herab, so daß immer ein Kompromiß geschlossen werden muß zwischen der optimalen Ausnutzung der Transistoreigenschaften und der

erforderlichen Arbeitspunktstabilisierung. Als höchste zulässige Temperatur überhaupt gelten folgende Werte:

Germanium-Transistoren:

75 °C (bis max. 90 °C)

Silizium-Transistoren:

150 °C (bis max. 200 °C).

Für jeden auch noch so einfachen Transistorverstärker sind entsprechende Schaltungsmaßnahmen zur Temperaturstabilisierung vorzusehen, es sei denn, der Transistor wird sehr weit unterhalb der zulässigen Leistungsgrenze betrieben.

4.4.5.1. Einstellen des Arbeitspunkts auf halbe Versorgungsspannung

Im Abschn. 2.1.3 haben wir bereits anhand eines einfachen Berechnungsbeispiels feststellen können, daß die infolge der Verlustleistung eintretende Erwärmung eines Halbleiter-Bauelements durch Vorschalten eines Widerstands herabgesetzt werden kann. Dabei ergab sich: Ist der Widerstandswert des Vorwiderstands genauso groß wie der Eigenwiderstand des Halbleiter-Bauelements, so wird die Verlustleistung am Halbleiter-Bauelement mit dessen zunehmender Erwärmung (Widerstandsverminderung) geringer. Da sich in einer Reihenschaltung von Widerständen die Teilspannungen wie die zugehörigen Widerstände verhalten, kann auch gesagt werden: Liegt in einer Reihenschaltung von ohmschem Widerstand und Halbleiter-Bauelement an beiden die gleiche Teilspannung — also je die halbe Gesamtspannung —, so nimmt die Verlustleistung am Halbleiter-Bauelement mit dessen Erwärmung ab.

Nun benötigen wir ja im Kollektorstromkreis ohnehin einen Verbraucherwiderstand (Lastwiderstand), den wir entsprechend bemessen können. Schaltungsmäßig sieht diese Maßnahme so aus:

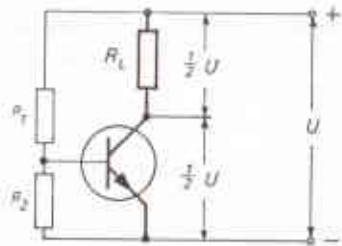


Abb. 4.36 — Temperaturstabilisierung durch Halbieren der Versorgungsspannung

Wegen der Festlegung auf halbe Versorgungsspannung kann natürlich auf optimale Spannungs-, Strom- oder Leistungsverstärkung keine Rücksicht mehr genommen werden. Daher wird dieses Verfahren vor allem für einfache, insbesondere Vorverstärker angewandt. Auf die Berechnung des Widerstandswerts für den Lastwiderstand R_L kommen wir noch zurück.

4.4.5.2. Einschalten eines Widerstands in die Emitterleitung (Stromgegenkopplung)

Liegt in der Emitterleitung eines Transistors ein ohmscher Widerstand R_E , so tritt an ihm im Betrieb ein Spannungsabfall auf. Dieser Spannungsabfall hängt von der Größe des Stroms I_E und des Widerstands R_E ab: $U_{RE} = I_E \cdot R_E$.

Da aber $I_E \approx I_C$ ist, kann man ohne großen Fehler auch schreiben

$$U_{RE} = I_C \cdot R_E.$$

In der Praxis wird R_E so gewählt, daß an ihm durch den Ruhestrom ein Spannungsabfall von $\boxed{U_{RE} \geq 0,1 \cdot U}$ erzeugt wird.

Je größer der Kollektorstrom bei bestimmtem Emitterwiderstand wird, desto größer wird der Spannungsabfall an R_E . Für die Schaltungsmaßnahme, die als Stromgegenkopplung bezeichnet wird, kann das folgende vereinfachte Schaltbild gezeichnet werden.

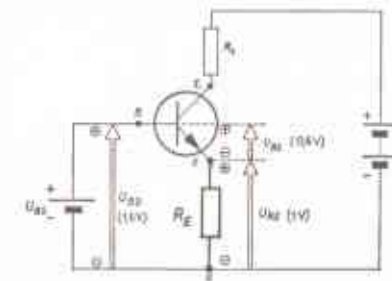


Abb. 4.37 — Temperaturstabilisierung durch Stromgegenkopplung (schematisch)

Zur Erklärung der Wirkungsweise dieser Schaltung müssen wir auf einige Grundlagenkenntnisse zurückgreifen. Nehmen wir einmal an, daß bei richtiger Arbeitspunkteinstellung am Widerstand R_E ein Spannungsabfall von 1 V auftritt (dieser Wert ergibt sich für $U = 9$ V). Das Potential am Emitter ist damit gegenüber dem Punkt 0 um 1 V positiver. Betrachtet man den Basis-Emitter-Stromkreis, so ergibt sich, daß der Spannungsabfall an R_E der angelegten Spannung U_{BE} entgegengerichtet ist. Da zum Betrieb des NPN-Transistors an seiner

Basis gegenüber dem Emitter ein um die Diffusionsspannung höheres Pluspotential liegen muß (bei Si-Transistoren 0,6 V), muß U_{B0} also $1\text{ V} + 0,6\text{ V} = 1,6\text{ V}$ betragen. Die Spannung U_{B0} wird mit Hilfe einer Spannungsteilerschaltung gewonnen, die bereits unter 4.4.3.3 besprochen worden ist. Damit ist die Spannung U_{BE} festgelegt. Steigen jetzt infolge Erwärmung des Transistors der Kollektorstrom und der Emitterstrom an, so tritt an R_E ein entsprechend größerer Spannungsabfall auf, z.B. 1,1 V. Die wirksame Basis-Emitter-Spannung beträgt dann nur noch

$$U_{BE} = U_{B0} - U_{RE} = 1,6\text{ V} - 1,1 = 0,5\text{ V};$$

sie ist also kleiner geworden. Wenn aber die Basis-Emitter-Spannung eines Transistors verringert wird, dann wird auch der Basisstrom und davon abhängig der Kollektorstrom kleiner. Die vollständige Schaltung ist in Abb. 4.38 dargestellt. Damit nicht infolge des Signalwechselstroms auch ein Wechselspannungsabfall an R_E auftritt, wird der Emitterwiderstand durch eine Kapazität überbrückt.

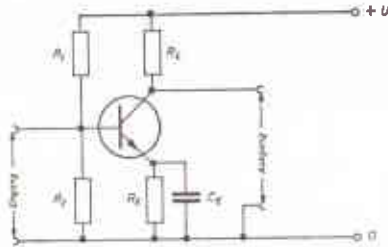


Abb. 4.38 — Transistorverstärkerstufe mit Stromgegenkopplung

Der **Vorteil** dieser Schaltungsmaßnahme gegenüber der Einstellung auf halbe Versorgungsspannung liegt darin, daß der Lastwiderstand R_L keinen unmittelbaren Einfluß hat und daher auf günstigste Spannungs-, Strom- oder Leistungsverstärkung bemessen werden kann. Als **Nachteil** dieser Schaltung ist anzugeben, daß die Versorgungsspannung um den an R_E im Ruhezustand auftretenden Spannungsabfall größer gewählt werden muß. An R_E wird eine zusätzliche Leistung verbraucht, die für den Verstärkungsvorgang nicht genutzt wird. Zudem liegt der Widerstand R_E des Basis-Spannungsteilers parallel zum Eingang des Transistors und stellt einen Nebenschluß für den Steuerwechselstrom dar. **Da die Wirkung dieser Schaltung von der Stärke des Kollektorstroms oder genauer des Emitterstroms abhängt, spricht man von der Stromgegenkopplung.**

4.4.5.3. Einschalten eines Widerstands zwischen Kollektor und Basis (Spannungsgegenkopplung)

Diese Schaltungsmaßnahme entspricht im Prinzip der unter 4.4.3.3 besprochenen Schaltung eines Basisvorwiderstands zur Erzeugung der erforderlichen Basis-Emitter-Spannung. Hier liegt der Basisvorwiderstand jedoch nicht direkt am Pluspol der Versorgungsspannung, sondern hinter dem Lastwiderstand — also am Kollektor.

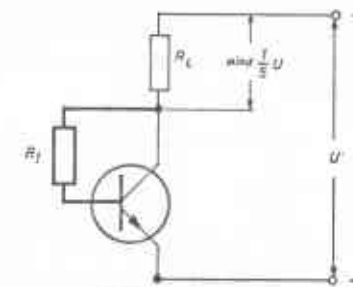


Abb. 4.39 — Temperaturstabilisierung durch Spannungsgegenkopplung

Die Wirkungsweise ist einfach: Steigt der Kollektorstrom infolge Erwärmung des Transistors, dann wird der Spannungsabfall am Lastwiderstand größer, die wirksame Kollektorspannung also kleiner. Da die Basis-Emitter-Spannung am Kollektor abgegriffen wird, verringert sie sich ebenfalls. Eine Verringerung der Basis-Emitter-Spannung hat aber bekanntlich eine Verringerung des Basisstroms und damit auch des Kollektorstroms zur Folge.

Vorteile dieser Stabilisierungsmaßnahme: Der Eingangswiderstand der Transistorstufe wird größer. Das ist z.B. wünschenswert, wenn die Steuerstromquelle selbst auch hochohmig ist (Anpassung!); z.B. haben Kristalltonabnehmer von Plattenspielern einen Widerstand von rd. $1\text{ M}\Omega$.

Nachteile: Die Schaltung ist aus Erfahrung nur dann wirksam, wenn an R_L mindestens $1/5$ der Versorgungsspannung abfällt. Zudem kann diese Gegenkopplung nicht wie die Stromgegenkopplung für Wechselstromvorgänge durch eine Kapazität unwirksam gemacht werden. **Die Wirkung der Schaltung ist von der Spannung am Kollektor abhängig; man spricht daher auch von der Spannungsgegenkopplung.** Neben der Stabilisierungswirkung werden zudem Signalverzerrungen herabgesetzt, die durch den Transistor selbst hervorgerufen werden — das geht allerdings auf Kosten des Verstärkungsfaktors.

4.4.5.4. Stabilisierung der Basis-Emitter-Spannung mit Hilfe eines Heißeleiters (NTC-Widerstands)

Wesentliches Kennzeichen eines Heißeleiters ist die Abnahme seines Widerstands mit zunehmender Temperatur. Der Widerstands-Temperaturbeiwert α ist daher — wie bei Halbleitern — negativ.

Schaltet man einen solchen Heißeleiter in einen Spannungsteiler parallel zu R_2 , so wirkt er temperaturstabilisierend: Mit zunehmender Temperatur wird der Widerstand des Heißeleiters kleiner. Dadurch fällt die Teilspannung und somit auch das Basis-Potential. Die Folge ist — wie in allen vorhergehenden Fällen — eine Verringerung des Basis- bzw. Kollektorstroms.

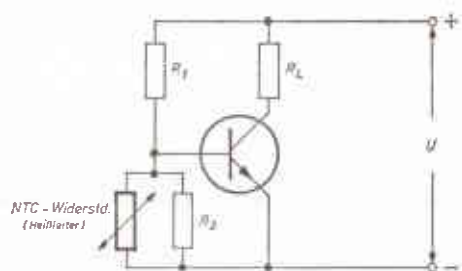


Abb. 4.40 — Temperaturstabilisierung durch NTC-Widerstand

Da die Kennlinie des NTC-Widerstands einen weitgehend anderen Verlauf aufweist als die Eingangskennlinie des Transistors, wird eine entsprechende Anpassung durch den Parallelwiderstand R_2 vorgenommen. Die Wirkung der Temperaturstabilisierung durch NTC-Widerstände hat aber im Gegensatz zu allen vorher besprochenen Schaltungsmaßnahmen eine ganz andere Ursache: Sie wird nicht durch die Zunahme des Kollektorstroms eines Transistors unmittelbar hervorgerufen, sondern durch die Umgebungstemperatur bzw. Gehäusetemperatur. Soll der NTC-Widerstand die Erwärmung des Transistors unmittelbar kompensieren, so muß er direkt am Transistorgehäuse oder auf das Kühlblech des Transistors montiert werden.

Vorteil: Es tritt kein Spannungs- und Leistungsverlust und keine Wechselstromgegenkopplung auf. Daher kann auch mit kleineren Versorgungsspannungen gearbeitet werden.

Nachteil: NTC-Widerstände sind teuer. Ihre Kennlinie muß fast ausnahmslos durch ohmsche Widerstände auf günstigste Wirkungsweise getrimmt werden.

4.4.5.5. Stabilisierung der Basis-Emitter-Spannung durch Dioden

Eine Diode ist in Durchlaßrichtung sehr hochohmig, wenn die anliegende Spannung kleiner als die Diffusionsspannung ist. Oberhalb der Diffusionsspannung nimmt ihr Widerstand ab. Da aber eine Diode ebenso wie ein Transistor aus einem Halbleiterstoff wie Germanium oder Silizium hergestellt ist, zeigt ihr Widerstandsverhalten auch eine starke Temperaturabhängigkeit: Mit steigender Temperatur verringert sich der Widerstand. Im Gegensatz zu einem NTC-Widerstand ist das Widerstands-Temperaturverhalten einer Diode dem eines Transistors sehr ähnlich. Darum haben sich Dioden zur Stabilisierung der Basis-Emitter-Spannung von Gegentakt-Transistorstufen einen festen Platz erobert.

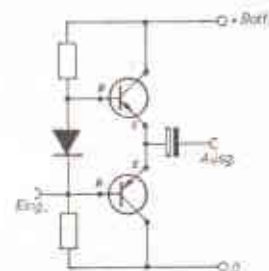


Abb. 4.41 — Stabilisierung der Basis-Emitter-Spannung einer Gegentaktstufe durch eine Diode

Leider haben Dioden infolge ihrer Begrenzerwirkung den Nachteil, daß sie z.B. Wechselspannungsspitzen abflachen und damit Verzerrungen verursachen. Sie können daher aus Gründen, auf deren Erläuterung in diesem Rahmen verzichtet werden muß, nur zur Stabilisierung der Basis-Emitter-Spannung in Transistor-Gegentaktstufen eingesetzt werden. In dem abgebildeten Schaltungsausschnitt liegt die Stabilisierungsdiode parallel zur Reihenschaltung der Basis-Emitter-Strecken der beiden Gegentakttransistoren. Ihre Diffusionsspannung muß gerade so groß sein wie die Summe der beiden Basis-Emitter-Spannungen.

Je nach Höhe der erforderlichen stabilisierten Spannung können höhere Spannungen durch die Reihenschaltung mehrerer Dioden, kleinere Spannungen durch zusätzliche Potentiometerschaltungen gewonnen werden. Die Stabilisierungsdioden müssen dabei grundsätzlich mit Vorwiderständen betrieben werden und stellen so den Teilwiderstand einer Spannungsteilerschaltung dar.

Vorteile: Geringer Schaltungsaufwand. Dioden bewirken eine gute Spannungsstabilisierung auch bei stark schwankender Versorgungsspannung. Kein Spannungs- bzw. Leistungsverlust und keine Gegenkopplung. Etwa gleiches Temperaturverhalten wie Transistoren; daher keine Kennlinienkorrektur wie bei NTC-Widerständen notwendig. Die Transistorstufe kann auf optimale Leistung ausgelegt werden.

Nachteil: Nur in Gegentaktstufen mit direkter Reihenschaltung der Basis-Emitter-Strecken möglich.

Zusammenfassend ist also festzuhalten: Da Transistoren in ihrer Wirkungsweise sehr stark temperaturabhängig sind, müssen in Transistorschaltungen zusätzliche Maßnahmen zur Temperaturstabilisierung vorgenommen werden.

Im einfachsten Fall stellt man den Arbeitspunkt auf halbe Versorgungsspannung ein. Soll aber die volle Leistung des Transistors ausgenutzt werden, so ist die Strom- oder Spannungsgegenkopplung anzuwenden. Für Transistorschaltungen, die großen Schwankungen der Umgebungstemperatur unterworfen sind (z.B. Meß- und Prüfgeräte für den Außendienst mit Sommer- und Winterbetrieb), wendet man vorwiegend die Temperaturstabilisierung durch NTC-Widerstände an. Bei Gegentaktstufen ist eine Arbeitspunktstabilisierung durch eine oder mehrere Dioden möglich.

4.4.6. Die richtige Bemessung des Lastwiderstands anhand des Ausgangskennlinienfeldes

Um die günstigsten Eigenschaften eines Transistorverstärkers zu erreichen oder überhaupt ein genaues Bild von der Arbeitsweise eines Transistorverstärkers zu erhalten, muß man sich vor allem des Ausgangskennlinienfeldes bedienen. Zum besseren Verständnis wollen wir hier etwas weiter ausholen, um die Zusammenhänge zu erläutern.

4.4.6.1. Die Reihenschaltung zweier ohmscher Widerstände

Im Abschn. 2.1.4 haben wir uns bereits mit den Kennlinien für ohmsche Widerstände beschäftigt. Dabei wurde das Verhalten eines oder mehrerer ohmscher Widerstände in einem Spannungs-Strom-Diagramm dargestellt. Jetzt soll einmal untersucht werden, wie der Kennlinienverlauf aussieht, wenn wir zwei Widerstände hintereinander, also in Reihe schalten. Dabei soll die Schaltung nach Abb. 4.42 zugrunde gelegt werden:

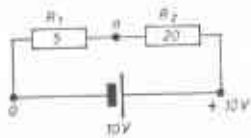


Abb. 4.42

Zur Darstellung der Kennlinien zweier in Reihe geschalteter Widerstände ist zu berücksichtigen, daß der Widerstand R_1 direkt am Potentialpunkt 0 und der Widerstand R_2 direkt am Potentialpunkt +10 V unserer Schaltung liegt. Daher muß die Kennlinie für R_1 im Nullpunkt beginnen, die Kennlinie R_2 im Punkt +10 V des Spannungs-Strom-Diagramms enden. Das Potential des Punkts a muß irgendwo zwischen 0 V und +10 V liegen. Diese Verhältnisse werden bei der Darstellung berücksichtigt: Wir stellen die Kennlinie des

Widerstands R_1 so dar, daß sie im Nullpunkt beginnt, die Kennlinie des Widerstands R_2 so, daß sie im Punkt +10 V endet. Da die Kennlinien für ohmsche Widerstände linear verlaufen, brauchen wir zu ihrer Darstellung keine ausführliche Wertetabelle aufzustellen. Es genügen schon zwei Punkte, um eine Gerade zu zeichnen. Formelmäßig ergibt sich folgender Zusammenhang für die Reihenschaltung zweier Widerstände:

$$U = I \cdot R_1 + I \cdot R_2$$

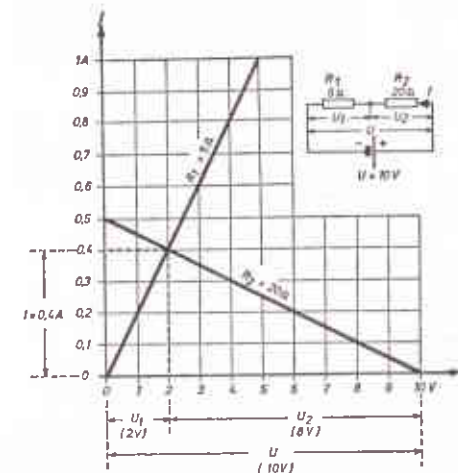
$$U = U_1 + U_2$$

$$\text{nach Umstellung } U_2 = U - U_1$$

Wir stellen eine einfache Wertetabelle auf:

U_1	0	5	10	V
$U_2 = U - U_1$	10	5	0	V
$I_1 = \frac{U_1}{R_1} = \frac{U_1}{5}$	0	1	2	A
$I_2 = \frac{U_2}{R_2} = \frac{U_2}{20}$	0,5	0,25	0	A

Anhand der Wertetabelle sind die Kennlinien im Spannungs-Strom-Diagramm darzustellen:



Im Diagramm ergibt sich ein Schnittpunkt für die beiden Widerstandskennlinien. In diesem Schnittpunkt ist demnach die Stromstärke in den beiden Widerständen R_1 und R_2 gleich groß. Das ist aber das Kennzeichen einer Reihenschaltung. Der Schnittpunkt stellt somit die Lösung unseres Problems dar. Wir können das leicht durch eine Rechnung beweisen:

Abb. 4.43 — Kennliniendarstellung für die Reihenschaltung zweier Widerstände

Gegeben: $U = 10 \text{ V}$; $R_1 = 5 \Omega$; $R_2 = 20 \Omega$

Gesucht: I ; U_1 ; U_2

$$\text{Lösung: } I = \frac{U}{R} = \frac{U}{R_1 + R_2} = \frac{10 \text{ V}}{5 \Omega + 20 \Omega} = \frac{10 \text{ V}}{25 \Omega} = 0,4 \text{ A}$$

$$U_1 = I \cdot R_1 = 0,4 \text{ A} \cdot 5 \Omega = 2 \text{ V}$$

$$U_2 = I \cdot R_2 = 0,4 \text{ A} \cdot 20 \Omega = 8 \text{ V}$$

Diese Werte ergeben sich genau aus dem Diagramm. Man könnte sich nun fragen, wozu dieser Aufwand gut sein soll, wenn die rechnerische Lösung viel einfacher und schneller ist. Doch zeigt bereits das nächste Beispiel, welchem Zweck die zeichnerische Lösung eines Reihenschaltungs-Problems dient.

4.4.6.2. Die Reihenschaltung einer Diode mit einem ohmschen Widerstand

Der Widerstand R_1 wird durch eine in Durchlaßrichtung geschaltete Diode ersetzt. Der Widerstand R_2 soll wie im ersten Beispiel $20\ \Omega$ und die Gesamtspannung 10 V betragen. Gesucht sind die Teilspannungen an der Diode und am Widerstand R_2 sowie die Stromstärke I der Reihenschaltung. Dieses Problem ist deswegen rechnerisch nicht mehr lösbar, weil es für den Kennlinienverlauf einer Diode keine mathematische Formel gibt.

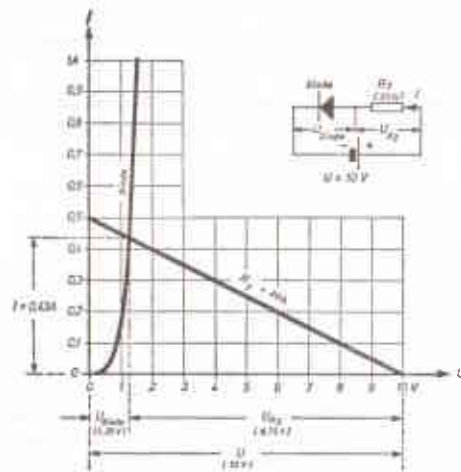


Abb. 4.44 — Kennliniendarstellung für die Reihenschaltung einer Diode mit einem ohmschen Widerstand R_2

Diese Lösung hätten wir durch Rechnung nicht finden können, da es — wie schon gesagt — für den Zusammenhang zwischen Strom und Spannung einer Diode keine Formel gibt.

4.4.6.3. Das Ermitteln des Vorwiderstands für eine Diode

An einem dritten Beispiel soll aufgezeigt werden, wie man ein Halbleiter-Bauelement durch geeignete Wahl des Vorwiderstands bis zur zulässigen Leistungsgrenze ausnutzen kann. Die Aufgabenstellung lautet: Gegeben ist die Kennlinie einer Diode. Die zulässige Verlustleistung der Diode beträgt 500 mW . Welchen Widerstandswert und welche Belastbarkeit muß ein Vorwiderstand mindestens haben, wenn die Gesamtspannung 5 V beträgt?

Hier muß also auf die in einem Meßversuch aufgenommene Diodenkennlinie zurückgegriffen werden. Da sich an der Gesamtspannung und dem Widerstand R_2 nichts geändert hat, können wir auf die Wertetabelle für R_2 des ersten Beispiels zurückgreifen und das Diagramm nach Abb. 4.44 darstellen.

Die Diodenkennlinie ist hier für eine Si-Diode dargestellt; aus dem Diagramm ergeben sich folgende Werte:

$$\begin{aligned} U_{\text{Diode}} &= 1,25\text{ V}, \\ U_{R2} &= 8,75\text{ V}, \\ I &= 0,43\text{ A}. \end{aligned}$$

Zur Lösung der Aufgabe müssen wir von der gegebenen Diodenkennlinie und der Leistungshyperbel für die zulässige Verlustleistung ausgehen. Die Kennlinie des Vorwiderstands muß jetzt so in das Diagramm gelegt werden, daß der Schnittpunkt mit der Diodenkennlinie gerade noch unter der Leistungshyperbel liegt. Da die Gesamtspannung mit 5 V gegeben ist, muß die Widerstandskennlinie bei $+5\text{ V}$ enden.

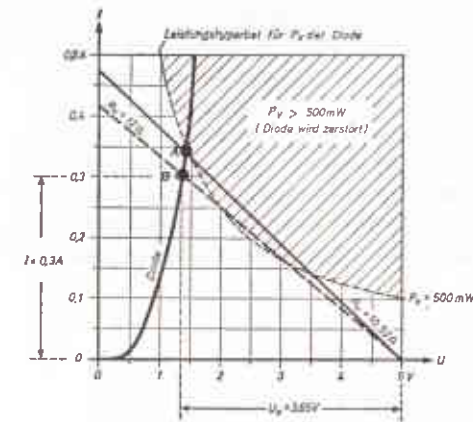


Abb. 4.45 — Ermittlung des Vorwiderstands R_v für eine Diode

Da es aber einen solchen Widerstand nach der Normreihe nicht gibt, wählen wir den nächstgrößeren Normwert $12\ \Omega$. Wird die Kennlinie des gewählten Widerstands $12\ \Omega$ in das Diagramm eingetragen, so ergibt sich ein Schnittpunkt B mit der Diodenkennlinie. Hierbei zeigt sich, was wir bereits wußten: Der Arbeitspunkt ist noch etwas weiter unter die Leistungshyperbel gesunken. Die Sicherheit vor einer Zerstörung der Diode ist größer geworden.

Die Belastung des Vorwiderstands kann aus der anliegenden Spannung und dem Betriebsstrom der Reihenschaltung ermittelt werden. Anhand des Diagramms ergibt sich für den $12\text{-}\Omega$ -Widerstand:

$$U_v = 3,65\text{ V}; \quad I = 0,3\text{ A}.$$

Daraus errechnet man eine Leistung von

$$P = U_v \cdot I = 3,65\text{ V} \cdot 0,3\text{ A} = 1,095\text{ W}$$

Der Vorwiderstand muß also mindestens mit $1,1\text{ W}$ belastbar sein.

4.4.6.4. Bestimmung des Lastwiderstands für einen Transistor

Da sich für die Kennlinien eines Transistors ebenso wenig wie für die der Dioden Formeln angeben lassen, wird zur Bestimmung des Lastwiderstands für einen Transistorverstärker genauso verfahren. Wir wollen das einmal anhand eines Ausgangskennlinienfeldes vornehmen. In Abb. 4.46 ist das Ausgangskennlinienfeld einschließlich Leistungshyperbel eines Transistors dargestellt. Die Versorgungs-

spannung soll mit 9 Volt (Batterie) festliegen. Gesucht wird der Widerstandswert für den Lastwiderstand im Kollektorstromkreis dieses Transistors für die **größtmögliche Leistung**.

Da das Ausgangskennlinienfeld aus vielen Einzelkennlinien besteht, muß die Widerstandskennlinie für den Lastwiderstand so liegen, daß **sämtliche** Schnittpunkte der Widerstandskennlinie mit den Ausgangskennlinien noch unter der Leistungshyperbel liegen. Diese Bedingung wird gerade noch erfüllt, wenn die Widerstandskennlinie die Leistungshyperbel berührt. Zur Festlegung des optimalen Lastwiderstands zeichnet man daher von der gegebenen Versorgungsspannung (hier 9 Volt) ausgehend eine Widerstandsgerade, die an der Leistungshyperbel liegt. Der Widerstandswert kann jetzt berechnet werden. Auf der Spannungsachse ergeben sich 9 V, auf der Stromachse 22,5 mA.

Der optimale Lastwiderstand beträgt also

$$R_{L1} = \frac{U}{I} = \frac{9 \text{ V}}{22,5 \text{ mA}} = \frac{9 \text{ V}}{0,0225 \text{ A}} = \underline{\underline{400 \, \Omega.}}$$

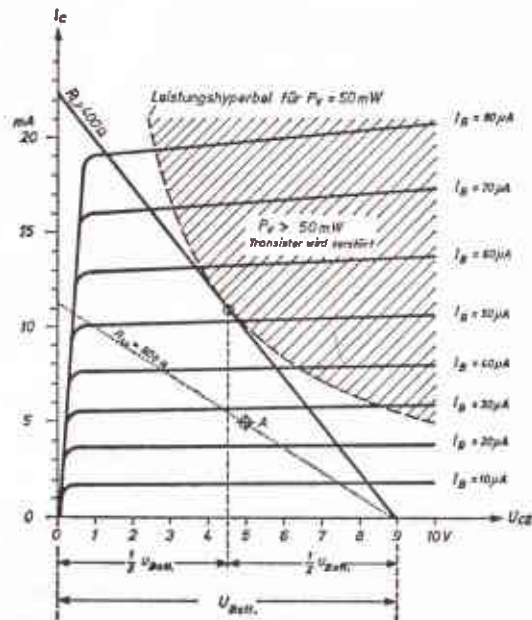


Abb. 4.46 — Ausgangskennlinienfeld für einen Transistor mit Widerstandsgeraden für optimalen Lastwiderstand und den Lastwiderstand für einen gewählten Arbeitspunkt A

In dem so ergänzten Kennlinienfeld lassen sich jetzt auch ein gewünschter Arbeitspunkt für den Transistor festlegen und die zugehörigen Werte bestimmen. Nehmen wir an, daß eine Temperaturstabilisierung durch Halbieren der Versorgungsspannung vorgenommen werden soll. Dann muß am Lastwiderstand und am Transistor je die halbe Gesamtspannung, also 4,5 V, liegen. Ziehen wir bei 4,5 V eine Senkrechte, so kann man den zugehörigen Arbeitspunkt bestimmen: Unter der Bedingung, daß am Transistor und am Lastwiderstand je die Hälfte der Gesamtspannung abfallen soll, liegt der Arbeitspunkt bei

$$I_B = 52 \, \mu\text{A} \text{ und}$$

$$I_C = 11 \text{ mA.}$$

Da der Schnittpunkt nicht genau auf der 50- μA -Kennlinie liegt, muß der Zwischenwert I_B geschätzt werden. Der Wert für I_C kann aber auch genau ermittelt werden, wenn die Steuerkennlinie $I_C = f(I_B)$ vorliegt. Dazu sucht man aus der Steuerkennlinie den zu $I_C = 11 \text{ mA}$ gehörenden Wert von I_B .

Etwas anders ist vorzugehen, wenn der Arbeitspunkt gegeben ist, z.B. durch die Angaben im Datenblatt des Transistors. Auch das wollen wir einmal üben. In dem dargestellten Ausgangs-Kennlinienfeld ist ein Arbeitspunkt A für $I_C = 5 \text{ mA}$ und $U_{CE} = 5 \text{ V}$ eingetragen. Dafür ergibt sich I_B zu rd. 27,5 μA .

Den genauen zu $I_C = 5 \text{ mA}$ gehörenden Wert für I_B entnimmt man der Strom-Steuerkennlinie $I_C = f(I_B)$. Zeichnet man die Widerstandsgerade des zum Arbeitspunkt A gehörenden Lastwiderstands, so schneidet sie die Stromachse bei 11,25 mA. Damit ergibt sich ein Lastwiderstand von

$$R_{L2} = \frac{U}{I} = \frac{9 \text{ V}}{11,25 \text{ mA}} = \frac{9 \text{ V}}{0,01125 \text{ A}} = \underline{\underline{800 \, \Omega.}}$$

4.4.7. Die Berechnung von Transistorverstärkern

In diesem Abschnitt soll aufgezeigt werden, wie Transistorverstärker anhand der Herstellerdaten verhältnismäßig einfach berechnet werden können. Für die Berechnungsbeispiele wird von folgenden Voraussetzungen ausgegangen:

Transistor BC 107

Versorgungsspannung 9 V

(zwei 4,5-V-Flachbatterien).

4.4.7.1. Berechnung des Basisvorwiderstands

Für einfache Vorverstärker mit Silizium-Transistoren, wie BC 107, genügt im allgemeinen die Erzeugung der Basis-Emitter-Spannung durch einen Basisvorwiderstand. Da nämlich Transistoren in Anfangsstufen meistens weit unterhalb der zulässigen Leistungsgrenze betrieben werden, ist hier eine Temperaturstabilisierung des Kollektorstroms nicht unbedingt notwendig. Den Arbeitspunkt entnimmt man am einfachsten den Datenblättern oder Tabellenbüchern. Dort sind für den betreffenden Transistor unter „Betriebsdaten“ bzw. „Kenn-daten“ die günstigsten Arbeitspunkte angegeben. Soll dagegen ein anderer als im Datenblatt angegebener Arbeitspunkt gewählt werden, so ist die richtige Bemessung der Bauelemente einer Transistor-schaltung nur mit Hilfe der Kennlinien möglich.

Bei Vorverstärkern ist noch zu beachten: Wird auf geringes Rauschen der Verstärkerstufe Wert gelegt, so darf der Kollektorstrom nicht größer als etwa 0,2 mA sein (z.B. Mikrofonverstärker). Sollen dagegen schon verhältnismäßig große Wechselspannungen oder -ströme (z.B. von Kristalltonabnehmern) verstärkt werden, dann wird der Kollektorstrom auf etwa 1 bis 2 mA eingestellt.

Im Datenblatt für den Transistor BC 107 findet man folgende Angaben über den Arbeitspunkt: $I_C = 2 \text{ mA}$ bei $U_{CE} = 5 \text{ V}$. Den folgenden Berechnungsbeispielen liegen die für einen bestimmten Transistor des Typs BC 107 gemessenen Werte (z.B. für I_B) zugrunde.

Beispiel 1 (Transistorkennlinien liegen vor):

Damit über den Transistor ein Kollektorstrom von 2 mA fließt, ist aufgrund der Messung des Gleichstromverstärkungsfaktors ein Basisstrom von $12 \mu\text{A}$ einzustellen. Aus der Eingangskennlinie $I_B = f(U_{BE})$ ermitteln wir den zugehörigen Spannungswert für U_{BE} . Er ergibt sich zu 0,64 V. Wir können jetzt ein Schaltbild mit den entsprechenden Spannungs- und Stromwerten zeichnen. Die bereits festliegenden Spannungs- und Stromwerte sind im Schaltbild eingerahmt.

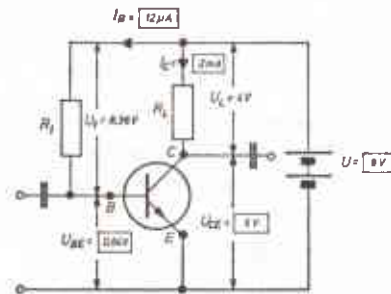


Abb. 4.47

Am Basisvorwiderstand R_1 muß demnach eine Spannung von

$$U_1 = U - U_{BE} = 9 \text{ V} - 0,64 \text{ V} = 8,36 \text{ V}$$

abfallen. Der Basisstrom wurde bereits mit $12 \mu\text{A}$ ermittelt. Aus Spannung und Strom ist mit Hilfe des Ohmschen Gesetzes der Widerstandswert zu berechnen:

$$R_1 = \frac{U_1}{I_B} = \frac{8,36 \text{ V}}{12 \cdot 10^{-6} \text{ A}} = \frac{8,36 \cdot 10^6}{12} \Omega = 0,697 \cdot 10^6 = \underline{\underline{697 \text{ k}\Omega}}$$

Der nächstliegende Normwert für Widerstände beträgt **680 k Ω** .

Beispiel 2 (Transistorkennlinien liegen nicht vor):

Für die Berechnung des Basisvorwiderstands ohne Kennlinien muß der Arbeitspunkt festliegen und der Gleichstromverstärkungsfaktor bekannt sein. Laut Datenblatt soll der Transistor BC 107 im Arbeitspunkt $I_C = 2 \text{ mA}$ und $U_{CE} = 5 \text{ V}$ einen Gleichstromverstärkungsfaktor B von 180 haben. Das ist — wie schon gesagt — natürlich ein Mittelwert.

Daraus können wir den Basisstrom berechnen:

$$\text{Aus } B = \frac{I_C}{I_B} \text{ wird nach Umstellung } I_B = \frac{I_C}{B} .$$

$$I_B = \frac{2 \text{ mA}}{180} = \frac{2000 \mu\text{A}}{180} = 11,1 \mu\text{A} .$$

Die Basis-Emitter-Spannung kann man für Silizium-Transistoren mit rd. 0,65 V annehmen. Somit ergibt sich für R_1 folgender Wert:

$$R_1 = \frac{U_1}{I_B}$$

$$U_1 = U - U_{BE} = 9 \text{ V} - 0,65 \text{ V} = 8,35 \text{ V}$$

$$R_1 = \frac{8,35 \text{ V}}{11,1 \cdot 10^{-6} \text{ A}} = \frac{8,35 \cdot 10^6}{11,1} \Omega = 0,753 \cdot 10^6 \Omega = \underline{\underline{753 \text{ k}\Omega}}$$

Nächstliegender Normwert **750 k Ω** .

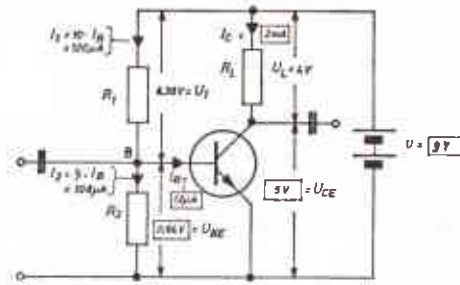
Dieses Beispiel läßt erkennen, daß die Berechnung anhand des Datenblatts zu einem anderen Ergebnis führt, als wenn man die genauen Meßwerte berücksichtigt.

4.4.7.2. Berechnung eines Basis-Spannungsteilers

Um mit Hilfe eines Basis-Spannungsteilers einen einigermaßen stabilen Arbeitspunkt für den Transistor zu errechnen, bemißt man ihn so, daß I_1 gleich dem zehnfachen Basisstrom ist. Der Einfluß der temperaturbedingten Widerstandsänderungen der Basis-Emitter-Strecke wird damit auf etwa 1/10 beschränkt. Es muß aber nochmals darauf hingewiesen werden, daß eine Arbeitspunktstabilisierung in Transistorstufen mit Basis-Spannungsteiler nur dann voll erreicht wird, wenn gleichzeitig Stromgegenkopplung angewandt wird.

Beispiel 1 (Transistorkennlinien liegen vor):

Der Basisstrom wurde für den Arbeitspunkt bereits durch Messung mit $12 \mu\text{A}$ ermittelt. Aus der Eingangskennlinie bestimmen wir wieder den zugehörigen Wert für U_{BE} , nämlich $0,64 \text{ V}$. Unsere Schaltung erhält jetzt folgendes Bild (wobei festliegende Werte eingerahmt sind):

**Abb. 4.48**

Über R_1 soll ein Strom I_1 fließen, der zehnmal so groß wie der Basisstrom I_B ist

$$I_1 = 10 \cdot I_B = 10 \cdot 12 \mu\text{A} = 120 \mu\text{A}.$$

Am Punkt B tritt eine Stromverzweigung auf

$$I_1 = I_B + I_2.$$

Somit fließt über R_2 noch die Differenz

$$I_2 = I_1 - I_B = 10 \cdot I_B - I_B = 9 \cdot I_B = 9 \cdot 12 \mu\text{A} = 108 \mu\text{A}.$$

Die Teilspannung an R_1 ergibt sich zu

$$U_1 = U - U_{BE} = 9 \text{ V} - 0,64 = 8,36 \text{ V}.$$

Damit können die Widerstände des Basis-Spannungsteilers berechnet werden:

$$R_1 = \frac{U_1}{I_1} = \frac{8,36 \text{ V}}{120 \cdot 10^{-6} \text{ A}} = \frac{8,36 \cdot 10^6}{120} \Omega = 0,0697 \cdot 10^6 \Omega = \underline{\underline{69,7 \text{ k}\Omega}}.$$

Nächster Normwert **68 k Ω** .

$$R_2 = \frac{U_{BE}}{I_2} = \frac{0,64 \text{ V}}{108 \cdot 10^{-6} \text{ A}} = \frac{0,64 \cdot 10^6}{108} \Omega = 0,00592 \cdot 10^6 \Omega = \underline{\underline{5,92 \text{ k}\Omega}}.$$

Nächster Normwert **6,2 k Ω** .

Beispiel 2 (Transistorkennlinien liegen nicht vor):

Hier dienen wieder die Angaben des Datenblatts als Berechnungsgrundlage:

$$I_C = 2 \text{ mA}; U_{CE} = 5 \text{ V}; \beta = 180.$$

$$I_B \text{ wird bestimmt aus } I_B = \frac{I_C}{\beta} = \frac{2000 \mu\text{A}}{180} = 11,1 \mu\text{A}.$$

Die Basis-Emitter-Spannung wird mit $0,65 \text{ V}$ angenommen. Damit wird

$$U_1 = U - U_{BE} = 9 \text{ V} - 0,65 \text{ V} = 8,35 \text{ V}.$$

Bei Bemessung des Spannungsteilers auf $I_1 = 10 \cdot I_B$ ergeben sich folgende Teilströme:

$$I_1 = 10 \cdot I_B = 10 \cdot 11,1 \mu\text{A} = 111 \mu\text{A}$$

$$I_2 = 9 \cdot I_B = 9 \cdot 11,1 \mu\text{A} = 100 \mu\text{A}.$$

Aus den zugehörigen Werten für Spannung und Strom können die Widerstände R_1 und R_2 des Basis-Spannungsteilers berechnet werden:

$$R_1 = \frac{U_1}{I_1} = \frac{8,35 \text{ V}}{111 \cdot 10^{-6} \text{ A}} = \frac{8,35 \cdot 10^6}{111} \Omega = 0,0752 \cdot 10^6 \Omega = \underline{\underline{75,2 \text{ k}\Omega}}.$$

Nächster Normwert **75 k Ω** .

$$R_2 = \frac{U_{BE}}{I_2} = \frac{0,65 \text{ V}}{100 \cdot 10^{-6} \text{ A}} = \frac{0,65 \cdot 10^6}{100} \Omega = 0,0065 \cdot 10^6 \Omega = \underline{\underline{6,5 \text{ k}\Omega}}.$$

Nächster Normwert **6,8 k Ω** .

Auch hier führt die überschlägliche Berechnung zu etwas anderen Ergebnissen als die Berechnung anhand der Meßwerte. Solange jedoch der Transistor-Arbeitspunkt weit unter der Leistungsgrenze liegt, sind diese Abweichungen nicht kritisch.

Da der Strom für den Spannungsteiler mit $10 \cdot I_B$ frei gewählt worden ist, können übrigens für den Spannungsteiler auch Widerstände mit den nächsthöheren oder niedrigeren Normwerten verwendet werden. Wichtig ist dabei nur, daß das Verhältnis von $R_1 : R_2$ gleichbleibt. Vergleichen wir unsere Ergebnisse von Beispiel 1 und 2, so beträgt das Verhältnis nach

$$\text{Beispiel 1: } \frac{R_1}{R_2} = \frac{68 \text{ k}\Omega}{6,2 \text{ k}\Omega} = \frac{11}{1}$$

$$\text{und nach Beispiel 2: } \frac{R_1}{R_2} = \frac{75 \text{ k}\Omega}{6,8 \text{ k}\Omega} = \frac{11}{1}$$

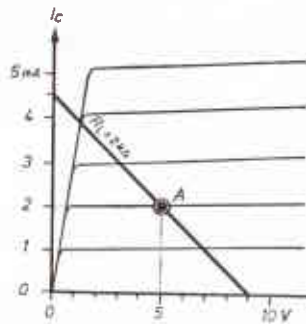
Der Transistorverstärker wird in beiden Fällen auf den gewünschten Arbeitspunkt eingestellt.

4.4.7.3. Bestimmung des Lastwiderstands R_L

Der für die höchstzulässige Leistung eines Transistors optimale Lastwiderstand läßt sich nur mit Hilfe des Ausgangskennlinienfeldes bestimmen (vgl. hierzu Abschn. 4.4.6.4). Für Transistoren in Vorverstärkerstufen mit einem Arbeitspunkt weit unterhalb der Leistungsgrenze kann man dagegen auch zur Bestimmung von R_L einfache Berechnungen anwenden.

Beispiel 1 (Transistorkennlinien liegen vor):

Der Arbeitspunkt $U_{CE} = 5 \text{ V}$ und $I_C = 2 \text{ mA}$ wird ins Ausgangskennlinienfeld eingezeichnet.

**Abb. 4.49**

Von der gegebenen Versorgungsspannung 9 V ausgehend zieht man eine Widerstandsgerade durch den Arbeitspunkt.

Die Widerstandsgerade schneidet die Stromachse bei $I_C \approx 4,5 \text{ mA}$.

Demnach muß der Lastwiderstand den Wert

$$R_L = \frac{U}{I} = \frac{9 \text{ V}}{4,5 \cdot 10^{-3} \text{ A}} = \underline{\underline{2000 \, \Omega}}$$

Beispiel 2 (Transistorkennlinien liegen nicht vor):

Wenn gemäß Datenblatt für den angegebenen Arbeitspunkt bei $I_C = 2 \text{ mA}$ am Transistor eine Spannung von $U_{CE} = 5 \text{ V}$ liegen soll, dann entfällt auf den Lastwiderstand die Teilspannung

$$U_L = U - U_{CE} = 9 \text{ V} - 5 \text{ V} = 4 \text{ V}.$$

Da wegen der Reihenschaltung über R_L der volle Kollektorstrom $I_C = 2 \text{ mA}$ fließt, läßt sich der Lastwiderstand einfach bestimmen:

$$R_L = \frac{U_L}{I_C} = \frac{4 \text{ V}}{2 \cdot 10^{-3} \text{ A}} = \underline{\underline{2000 \, \Omega}}.$$

4.4.7.4. Berechnung der Kopplungskapazitäten

Für Vorverstärker wird nahezu ausschließlich die C-Kopplung angewendet. Das Berechnungsbeispiel bezieht sich auf einen Übertragungsbereich von $f_u = 30 \text{ Hz}$ bis $f_o = 15000 \text{ Hz}$.

Der Eingangswiderstand der Transistorstufe muß aus der Parallelschaltung von R_2 des Basis-Spannungsteilers mit dem Basis-Emitter-Bahnwiderstand R_{BE} des Transistors bestimmt werden. R_2 liegt mit $6,8 \text{ k}\Omega$ fest. Am Ausgang soll ein Kopfhörer mit $R_K = 4000 \, \Omega$ liegen.

$$R_{\text{ein}} = \frac{R_2 \cdot R_{BE}}{R_2 + R_{BE}}$$

$$R_{BE} = \frac{U_{BE}}{I_B} = \frac{0,64 \text{ V}}{12 \cdot 10^{-6} \text{ A}} = \frac{0,64 \cdot 10^6}{12} \, \Omega$$

$$R_{BE} = 0,0533 \cdot 10^6 \, \Omega = 53,3 \text{ k}\Omega$$

$$R_{\text{ein}} = \frac{6,8 \text{ k}\Omega \cdot 53,3 \text{ k}\Omega}{6,8 \text{ k}\Omega + 53,3 \text{ k}\Omega} = \frac{362}{60,1} \text{ k}\Omega \approx \underline{\underline{6 \text{ k}\Omega}}.$$

Die Kopplungskapazitäten werden mit Hilfe der bereits bekannten Formel berechnet:

$$C_{\text{Kein}} = \frac{1}{2 \pi \cdot f_u \cdot R_{\text{ein}}} = \frac{1}{2 \cdot 3,14 \cdot 30 \cdot 6000} \text{ F}$$

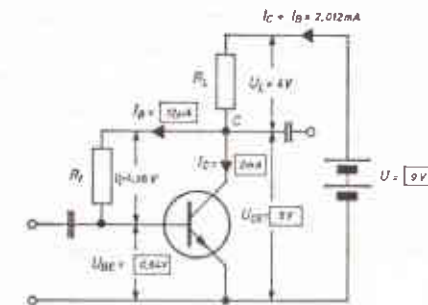
$$C_{\text{Kein}} = 0,000000885 \text{ F} = 0,885 \, \mu\text{F} \approx \underline{\underline{1 \, \mu\text{F}}}, \text{ gewählt } 4,7 \, \mu\text{F}.$$

$$C_{\text{Kaus}} = \frac{1}{2 \pi \cdot f_o \cdot R_K} = \frac{1}{2 \cdot 3,14 \cdot 30 \cdot 4000} = 0,0000013 \text{ F}$$

$$= 1,3 \, \mu\text{F} \approx \underline{\underline{1,5 \, \mu\text{F}}}, \text{ gewählt } 6,8 \, \mu\text{F}.$$

4.4.7.5. Vorverstärker mit Spannungsgegenkopplung

Hier soll an einem Beispiel die Bestimmung der entsprechenden Widerstandswerte aufgezeigt werden. Auch für dieses Beispiel soll der gleiche Arbeitspunkt $U_{CE} = 5 \text{ V}$ und $I_C = 2 \text{ mA}$ des Transistors BC 107 gelten. Anhand der durch den gegebenen Arbeitspunkt festliegenden Werte kann das Schaltbild nach Abb. 4.50 gezeichnet werden (wobei festliegende Werte wieder eingerahmt sind):

**Abb. 4.50**

Bei Spannungsgegenkopplung — die die im Transistor auftretenden Verzerrungen weitgehend kompensiert und eine Temperaturstabilisierung bewirkt — wird die Basis-Emitter-Spannung über den Widerstand R_1 am Punkt C abgegriffen. Da R_1 als Basisvorwiderstand geschaltet ist, fließt über ihn der Basisstrom $I_B = 12 \, \mu\text{A}$ (Meßwert!).

An R_1 muß eine Spannung von

$$U_1 = U_{CE} - U_{BE} = 5 \text{ V} - 0,64 \text{ V} = 4,36 \text{ V abfallen.}$$

Der Widerstand ergibt sich zu

$$R_1 = \frac{U_1}{I_B} = \frac{4,36 \text{ V}}{12 \cdot 10^{-6} \text{ A}} = 0,363 \cdot 10^6 \, \Omega = \underline{\underline{363 \text{ k}\Omega}}.$$

Nächster Normwert **360 kΩ**.

Der Eingangswiderstand R_{ein} der Transistorstufe wird hier ausschließlich durch den Basis-Emitter-Bahnwiderstand R_{BE} des Transistors gebildet.

Er beträgt

$$R_{\text{BE}} = \frac{U_{\text{BE}}}{I_{\text{B}}} = \frac{0,64 \text{ V}}{12 \cdot 10^{-6} \text{ A}} = \underline{\underline{53,3 \text{ k}\Omega}}$$

und ist somit nahezu 10mal so groß wie bei einer Transistorstufe mit Basis-Spannungsteiler.

Bei der Berechnung des Lastwiderstands R_{L} ist zu berücksichtigen, daß über ihn zusätzlich auch der Basisstrom fließt (siehe Schaltung). Da an R_{L} eine Spannung von 4 V liegt, ist die Bedingung $U_{\text{L}} \geq \frac{1}{5} U$ für die Wirksamkeit der Spannungsgegenkopplung erfüllt.

$$R_{\text{L}} = \frac{U_{\text{L}}}{I_{\text{C}} + I_{\text{B}}} = \frac{4 \text{ V}}{2,012 \cdot 10^{-3} \text{ A}} = 1999 \Omega \approx \underline{\underline{2 \text{ k}\Omega}} \text{ (Normwert).}$$

4.4.7.6. Transistorverstärker mit Stromgegenkopplung

Stromgegenkopplung wird zur Temperaturstabilisierung des Arbeitspunkts am häufigsten angewandt. Die Berechnung der Widerstandswerte ist zwar etwas aufwendiger, bereitet jedoch bei Beherrschung der Gesetzmäßigkeiten der Gleichstromlehre keine Schwierigkeiten. Kennzeichen der Stromgegenkopplung ist der Widerstand R_{E} in der Emittierzuleitung. Für $U = 9 \text{ V}$ wird R_{E} im allgemeinen so bemessen, daß an ihm ein Spannungsabfall von 1 Volt auftritt. Für den gegebenen Arbeitspunkt lassen sich in der Verstärkerschaltung folgende Spannungs- und Stromwerte angeben (bereits festliegende Werte sind eingekreist):

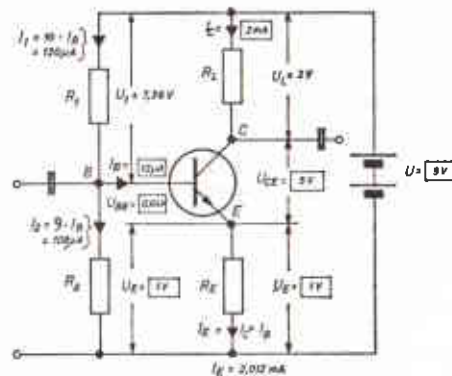


Abb. 4.51

Zur Berechnung der Widerstände R_1 , R_2 , R_{L} und R_{E} müssen jeweils anliegende Spannung (Spannungsabfall) und durchfließender Strom bekannt sein. Die Größen ergeben sich wie folgt.

Für R_1 :

$$U_1 = U - (U_{\text{E}} + U_{\text{BE}}) = 9 \text{ V} - (1 \text{ V} + 0,64 \text{ V}) = 7,36 \text{ V}$$

$$I_1 = I_0 + I_{\text{B}} \text{ (Richtwert)} = 10 + 12 \mu\text{A} = 120 \mu\text{A}$$

$$R_1 = \frac{U_1}{I_1} = \frac{7,36 \text{ V}}{120 \cdot 10^{-6} \text{ A}} = 0,0613 \cdot 10^6 \Omega = 61,3 \cdot 10^3 \Omega = \underline{\underline{61,3 \text{ k}\Omega}}.$$

Nächster Normwert **62 kΩ**.

Für R_2 :

$$U_2 = U_{\text{E}} + U_{\text{BE}} = 1 \text{ V} + 0,64 \text{ V} = 1,64 \text{ V}$$

(U_2 wird auch als U_{B0} bezeichnet.)

$$I_2 = I_0 + I_{\text{B}} = 9 + 12 \mu\text{A} = 108 \mu\text{A}$$

$$R_2 = \frac{U_2}{I_2} = \frac{1,64 \text{ V}}{108 \cdot 10^{-6} \text{ A}} = 0,0152 \cdot 10^6 \Omega = 15,2 \cdot 10^3 \Omega = \underline{\underline{15,2 \text{ k}\Omega}}.$$

Nächster Normwert **15 kΩ**.

Für R_{L} :

$$U_{\text{L}} = U - (U_{\text{CE}} + U_{\text{E}}) = 9 \text{ V} - (5 \text{ V} + 1 \text{ V}) = 3 \text{ V}$$

(Sie sehen, daß hier U_{L} kleiner geworden ist. U_{E} vermindert die ausnutzbare Betriebsspannung!)

$$I_{\text{L}} = I_{\text{C}} = 2 \text{ mA}$$

$$R_{\text{L}} = \frac{U_{\text{L}}}{I_{\text{L}}} = \frac{3 \text{ V}}{2 \cdot 10^{-3} \text{ A}} = 1,5 \cdot 10^3 \Omega = \underline{\underline{1,5 \text{ k}\Omega}} \text{ (Normwert).}$$

Für R_{E} :

$$U_{\text{E}} = 1 \text{ V}$$

$$I_{\text{E}} = I_{\text{C}} + I_{\text{B}} = 2 \text{ mA} + 12 \mu\text{A} = 2000 \mu\text{A} + 12 \mu\text{A} = 2012 \mu\text{A} = 2,012 \text{ mA}$$

$$R_{\text{E}} = \frac{U_{\text{E}}}{I_{\text{E}}} = \frac{1 \text{ V}}{2,012 \cdot 10^{-3} \text{ A}} = 0,497 \cdot 10^3 \Omega = \underline{\underline{497 \Omega}}.$$

Nächster Normwert **510 Ω**.

Zur Berechnung des Überbrückungskondensators für $f_{\text{u}} = 30 \text{ Hz}$ bedienen wir uns wieder der für C-Kopplung angegebenen Formel.

$$C_{\text{E}} = \frac{1}{2 \cdot \pi \cdot f_{\text{u}} \cdot R_{\text{E}}} = \frac{1}{2 \cdot 3,14 \cdot 30 \cdot 510} = 0,0000104 \text{ F} = \underline{\underline{10,4 \mu\text{F}}}, \text{ gewählt } 47 \mu\text{F}.$$

Die aufgeführten Berechnungsbeispiele sind bewußt einfach gehalten und nur für Vor- und Treiberstufen anwendbar. Leistungsstufen werden im allgemeinen als Gegentaktstufen mit sehr niedrig liegendem Arbeitspunkt ausgelegt (sog. B-Verstärker). Auf die Berechnung sol-

cher Schaltungen sowie auf die genaue Bestimmung des Temperaturverhaltens von Transistorstufen muß im Rahmen dieses Lehrbuches verzichtet werden. Um die Berechnung einfacher Vor- und Treiberstufen zu erleichtern, sollen die wichtigsten Punkte noch einmal kurz zusammengestellt werden:

a) Auswahl der Schaltungsart

Vorverstärker: Basisvorwiderstand oder Basis-Spannungsteiler. Basisvorwiderstand als Spannungsgegenkopplung ergibt hohen, Basis-Spannungsteiler niedrigen Eingangswiderstand (Anpassung an die Steuerstromquelle beachten!).

Treiber- oder Endstufen: Basis-Spannungsteiler und ausreichend großer Emitterwiderstand. Sollen in der Verstärkerstufe auftretende Verzerrungen vermieden werden, so kann der Basis-Spannungsteiler am Kollektor abgegriffen werden (Spannungsgegenkopplung).

b) Festlegung der Versorgungsspannung

Vorverstärker: Es genügen verhältnismäßig kleine Spannungen von ca. 3 bis 12 V.
Treiber- oder Endstufen: Mit höheren Versorgungsspannungen lassen sich größere Verstärkerleistungen oder zumindest höhere Spannungsverstärkungen erzielen. U etwa 6 bis 50 V.

c) Wahl des Transistors und Ermittlung des Arbeitspunkts

Vorverstärker: Allgemein werden dafür Kleisleistungstransistoren mit P_V um etwa 100 mW verwendet. Für rauscharme Verstärker gibt es besondere, rauscharme Typen.

Treiber- oder Endstufen: Je nach Leistungsbedarf Transistoren mit höherer Verlustleistung.

Arbeitspunkt: Liegen keine Kennlinien vor, so entnimmt man den Arbeitspunkt der Spalte „Kenndaten“ des Datenblatts. Für rauscharme Transistoren soll I_C höchstens 0,2 bis 0,3 mA betragen. Als Basis-Emitter-Spannung U_{BE} wird für Vorverstärker mit Ge-Transistoren 0,15 bis 0,3 V und mit Si-Transistoren 0,55 bis 0,65 V gewählt. Für Treiber- und Endstufen ist U_{BE} etwas höher zu wählen (Ge ca. 0,3 bis 0,38 V und Si ca. 0,66 bis 0,72 V). Grundsätzlich gehört zu einem niedrigen Arbeitspunkt ein niedriger Wert für U_{BE} .

d) Ermittlung des Verstärkungsfaktors B

Der Verstärkungsfaktor B wird zweckmäßig durch Messung ermittelt, da auch bei Transistoren gleichen Typs weite Streuungen für B auftreten.

e) Berechnung von I_B für den Arbeitspunkt

Mit Hilfe der Formel $I_B = \frac{I_C}{B}$ durch Einsetzen der Werte für I_C (Datenblatt) und B (Meßwert).

f) Zeichnen der Schaltung mit Spannungs- und Stromwerten

Die Darstellung einer ausführlichen Schaltung mit allen Spannungs- und Stromwerten ist besonders wichtig und erleichtert die Berechnung der Widerstandswerte sehr. Einzelne Spannungs- und Stromwerte anhand unserer Berechnungsbeispiele ermitteln.

g) Berechnung der Widerstandswerte

Liegen die Spannungs- und Stromwerte fest, so ist die Berechnung der Widerstandswerte mit Hilfe des Ohmschen Gesetzes sehr einfach.

h) Berechnung der Kopplungs- und Überbrückungskapazitäten

Sie wird für die niedrigste noch zu verstärkende Frequenz nach der Formel

$C_k = \frac{1}{2 \cdot \pi \cdot I_u \cdot R}$ vorgenommen, wobei R der gesamte hinter dem Kondensator liegende Widerstand ist. In der Praxis wird der etwa 3- bis 5fache Kapazitätswert gewählt.

4.4.8. Der praktische Aufbau eines Vorverstärkers

Nach der Berechnung der benötigten Bauteile können wir jetzt an den Aufbau des Verstärkers gehen. Einfachste und billigste Mittel zum Aufbau eines kleinen Verstärkers sind Lötösenstreifen oder Pertinaxplatten von 1 bis 2 mm Dicke. Für größere Verstärker lohnt sich dagegen auch die Herstellung „gedruckter“ Schaltungen aus Pertinaxplatten, die mit Kupferbahnen oder Kupferscheiben kaschiert und in bestimmten Lochabständen gebohrt sind. Unser erster selbst berechneter und hergestellter Verstärker soll ein Vorverstärker mit Basisvorwiderstand als Spannungsgegenkopplung sein. Die entsprechenden Widerstandswerte haben wir bereits für einen Transistor BC 107 mit dem Arbeitspunkt $U_{CE} = 5 \text{ V}$ bzw. $I_C = 2 \text{ mA}$ berechnet. In unsere Verstärkerschaltung können wir also bereits feste Werte eintragen.

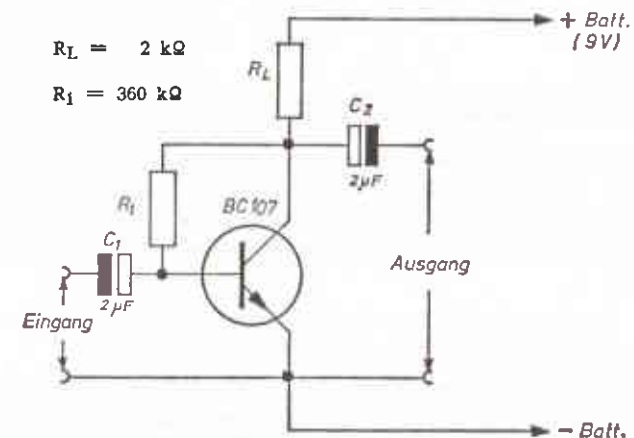


Abb. 4.52 — Schaltung des Vorverstärkers

Dieser Verstärker eignet sich z.B. vorzüglich als Vorverstärker für Kristalltonabnehmer oder hochohmige dynamische Mikrofone oder auch für Diodenempfänger. Er reicht aus, um die Wechselströme eines Kristalltonabnehmers oder eines Diodenempfängers in einem angeschlossenen Kopfhörer ($4000\ \Omega$) gut hörbar zu machen. Er kann aber auch als Mikrofonvorverstärker für einen nachgeschalteten Leistungsverstärker benutzt werden. Die Verstärkung vieler Leistungsverstärker reicht nicht zum Betrieb eines dynamischen Mikrofons aus.

4.4.8.1. Aufbau auf Lötösenstreifen

Die Bauteile sollten so angeordnet werden, daß Eingang und Ausgang des Verstärkers möglichst jeweils am Ende des Lötösenstreifens liegen. Das erleichtert den Anschluß der Steuerstromquelle und des Verbrauchers (z.B. Kopfhörer). Unser Verstärker würde — auf einen Lötösenstreifen montiert — etwa so aussehen:

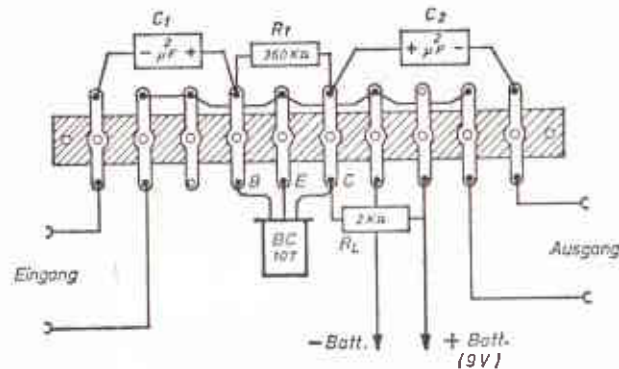


Abb. 4.53 — Aufbau des Vorverstärkers (Lötösenstreifen)

Dieses Verfahren hat den Vorteil, daß man zum Befestigen der Bauteile keine Bohrungen anzufertigen braucht. Bei Bedarf kann die Schaltung leicht geändert und der Lötösenstreifen ohne Schwierigkeiten montiert werden. Unter allen Umständen ist aber darauf zu achten, daß

- die Anschlußdrähte der Bauteile nicht unmittelbar an der Austrittsstelle gebogen werden und
- Halbleiter-Bauelemente, wie Transistoren, beim Lötten nicht zu warm werden. Sie können sonst zerstört werden. Schutzmaßnahmen: Entweder Anschlußdrähte lang lassen oder die Anschlußdrähte beim Anlöten mit einer Flachzange kühlen.

4.4.8.2. Aufbau auf Pertinaxplatte

Auch hier ist ein sorgfältiger Entwurf wichtig. Die Bauteile werden so angeordnet, daß man mit möglichst geraden und kurzen Verbindungen auskommt. Die Anschlüsse für Ein- und Ausgang sollten nicht irgendwo aus der Schaltung abgegriffen werden.

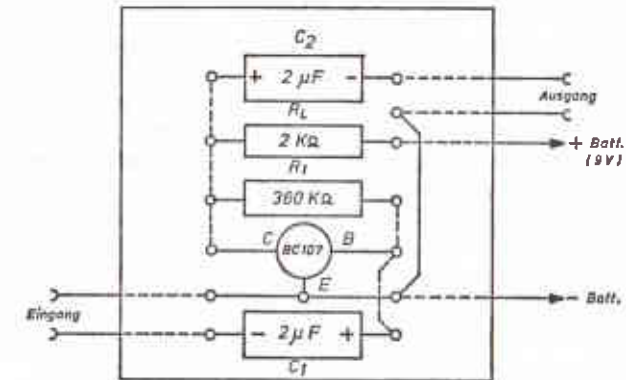


Abb. 4.54 — Aufbau des Vorverstärkers (Pertinaxplatte)

Unter der Platte liegende Verbindungen sind gestrichelt gezeichnet. Anhand der Aufbauskizze werden die Bauteile auf die Pertinaxplatte gelegt und das Abstandsmaß der Bohrungen ermittelt. Die Anschlußdrähte führt man von oben durch die Bohrung und biegt auf der Unterseite mit Hilfe einer Spitzzange Ösen. An diese Ösen können die Verbindungsdrähte leicht angelötet werden.

Alle für den Aufbau eines Transistorverstärkers vorgesehenen Bauteile sollten vor dem Einbau geprüft und gemessen werden; hierbei handelt es sich um folgende Aufgaben:

- Prüfung des Transistors,
- Messung des genauen Widerstandswerts und ggf. Auswahl eines geeigneten Widerstands,
- Prüfung der Kondensatoren auf Isolationswiderstand. Sie wird mit dem höchsten Widerstandsbereich eines Ohmmeters vorgenommen. Die Widerstandsmessung muß für Kondensatoren $\infty\ \Omega$ ergeben. Dabei ist jedoch zu berücksichtigen, daß sich Kondensatoren zunächst aufladen. Das hat einen Zeigerausschlag zur Folge, der jedoch je nach Größe der Kapazität schnell oder langsam auf den Wert ∞ zurückgeht. Bei Elektrolytkondensatoren größerer Kapazität ergeben sich dagegen häufig Widerstände von einigen 100 Kiloohm.

Beide Verfahren zum Aufbau eines Transistorverstärkers bieten die Möglichkeit, ganz nach Bedarf eine Reihe von Verstärkerstufen einzeln aufzubauen und hintereinanderschalten. Sowohl die Lötösenstreifen als auch die Pertinaxplatten lassen sich parallel nebeneinan-

der auf einem Rahmen befestigen. Man hat damit die Möglichkeit, einen Verstärker weiter auszubauen oder auch einzelne Stufen zu ändern („D-Zug-Verstärker“). Gegenüber gedruckten oder geätzten Schaltungen haben die beschriebenen einfachen Aufbauarten noch den Vorteil, daß sie bei Bedarf leicht geändert werden können.

4.4.9. Mehrstufige Transistorverstärker

Die von üblichen Signalquellen — wie Mikrofonen, Tonabnehmern oder Diodenausgängen von Rundfunkempfängern — abgegebene Leistung von ca. 0,01 bis 1 μW reicht nicht einmal aus, einen empfindlichen Kopfhörer ausreichend laut zu betreiben. Zur Lautsprecherwiedergabe von Sprache und Musik sind daher mehrstufige Verstärker notwendig. Ein Lautsprecher benötigt für eine Wiedergabe mit Zimmerlautstärke ca. 0,2 bis 0,5 Watt, bei größerem Lautstärkebedarf ca. 1 bis 3 Watt und zur Beschallung größerer Räume zwischen 10 und 1000 Watt elektrische Leistung. Um also beispielsweise das Signal eines Tonabnehmers für Schallplatten mit 0,1 μW in einem Lautsprecher mit Zimmerlautstärke wiederzugeben, ist eine Leistungsverstärkung von

$$v_p = \frac{P_2}{P_1} = \frac{0,5 \text{ W}}{0,1 \mu\text{W}} = \frac{500 \text{ mW}}{0,0001 \text{ mW}} = 5\,000\,000,$$

also fünfmillionenfach erforderlich. Ein so hoher Verstärkungsfaktor ist nur mit mehrstufigen Verstärkern erreichbar. Bei mehrstufigen Verstärkern ist zwischen drei Stufen zu unterscheiden:

Vorverstärkerstufe: Sie hat die Aufgabe, die im allgemeinen sehr kleinen Spannungen der Signalquellen (Mikrofone, Tonabnehmer usw.) zu verstärken und soll so ausgelegt sein, daß ihr Eingangswiderstand an den Innenwiderstand der Signalquelle angepaßt ist. Je nach Verstärkungsbedarf kann sie aus einem oder mehreren Transistoren bestehen. Ggf. wird in oder hinter der Vorverstärkerstufe eine Klangregelung vorgenommen. Die Vorverstärkerstufe darf nur geringes Rauschen verursachen (rauscharme Transistoren mit niedrigem Arbeitspunkt).

Treiberstufe: In ihr findet der Übergang von der sogenannten Kleinsignalverstärkung zur Leistungsverstärkung statt. Sie treibt die Leistungsverstärkerstufe. Da ein Leistungsverstärker zur Steuerung schon eine bestimmte Leistung benötigt, wird die Treiberstufe auf optimale Leistungsverstärkung ausgelegt. Die Treiberstufe ist zudem Anpassungsglied zwischen Vorverstärker- und Leistungsverstärkerstufe.

Leistungsverstärkerstufe: Ihre Aufgabe ist es, die zum Betrieb des Verbrauchers (z.B. Lautsprecher) benötigte Leistung zu liefern. Zur

Erzielung eines niedrigen Kollektorruhestroms werden Leistungsverstärkerstufen nahezu ausschließlich als Gegentaktstufen gebaut. Vereinfacht kann die Wirkungsweise einer Gegentaktstufe so beschrieben werden: Der Arbeitspunkt der beiden Gegentakttransistoren wird sehr niedrig eingestellt. Damit fließt ohne Wechselstromsignal ein verhältnismäßig geringer Kollektorruhestrom. Man spricht bei dieser Art der Arbeitspunkteinstellung auch von B-Verstärkung. Beim Anlegen eines Wechselstromsignals wird einer der Gegentakttransistoren nur von der positiven, der andere nur von der negativen Halbwelle des Wechselstroms durchgesteuert. In den Transistoren fließt also nur je während einer Halbwelle des Steuerwechselstroms ein höherer Kollektorstrom. Den schaltungstechnisch geringsten Aufwand erfordern Gegentaktstufen mit Komplementärtransistoren. Ein Komplementärpaar besteht aus einem NPN- und einem PNP-Transistor. Die zu einem Paar gehörenden Transistoren müssen gleiche Eigenschaften haben und daher sorgfältig ausgesucht werden.

4.4.9.1. Schaltung eines vollständigen Verstärkers mit 10 W Ausgangsleistung

Der hier beschriebene Verstärker wird auch gehobenen Ansprüchen gerecht und besteht aus drei Baugruppen:

- a) **Regelteil** mit Lautstärke-, Tiefen- und Höhenregler,
- b) **Verstärkerteil** mit Vorverstärker-, Treiber- und Leistungsverstärkerstufe sowie
- c) **Stromversorgungsstell** (Netzgerät).

Im Eingang des Regelteils liegt ein Potentiometer zur Lautstärkeeinstellung. Darauf folgt eine Verstärkerstufe mit dem besonders rauscharmen Si-Transistor BC 109.

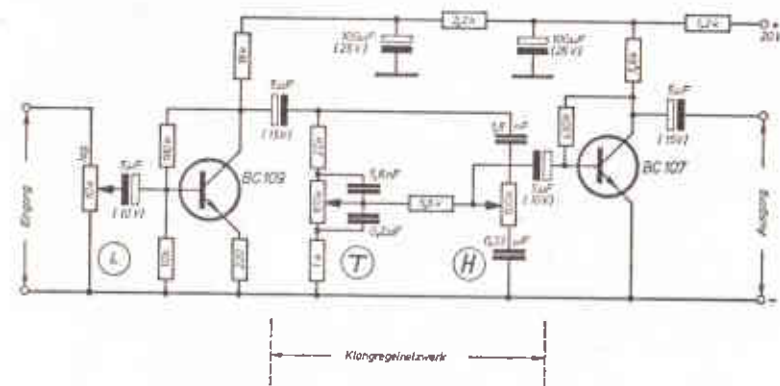


Abb 4.55 — Schaltung des Regelteils

Die Stufe ist spannungs- und stromgegekoppelt, um auch bei hohem Eingangssignal Verzerrungen zu vermeiden. Sie soll das Eingangssignal verstärken, da im nachfolgenden Klangregelnetzwerk verhältnismäßig hohe Verluste auftreten. Das eigentliche Klangregelnetzwerk ermöglicht die Beeinflussung des Frequenzgangs in weiten Bereichen. Steht der Schleiferabgriff des Tiefenreglers T oben, so werden Tonfrequenzen von etwa 40 Hz um 18 dB (das ist auf das Achtfache) angehoben. In der unteren Stellung werden die tiefen Frequenzen um 12 dB abgesenkt (etwa auf ein Viertel des mittleren Werts). Die obere Stellung des Höhenreglers H bedeutet eine Anhebung der hohen Frequenzen bei etwa 16 000 Hz um 18 dB (auf das Achtfache), während sie in der unteren Stellung um 12 dB (auf ein Viertel) gesenkt werden. Zur Entkopplung zwischen dem Klangregelnetzwerk und der Vorverstärkerstufe des eigentlichen Verstärkerteils dient die Verstärkerstufe mit dem Transistor BC 107.

Der Vorverstärker im Verstärkerteil enthält einen PNP-Transistor BC 177. Wegen der Schichtfolge PNP muß dieser Transistor mit dem Emittor an Plus (nach „oben“) in die Schaltung eingefügt werden. Er erhält über die Emittorzuleitung (1,2 k Ω) eine kräftige Stromgegenkopplung aus der Leistungsverstärkerstufe. Dadurch werden Verzerrungen stark herabgesetzt. Das RC-Glied 39 Ω — 250 μ F dient zur Frequenzkorrektur der Gegenkopplung.

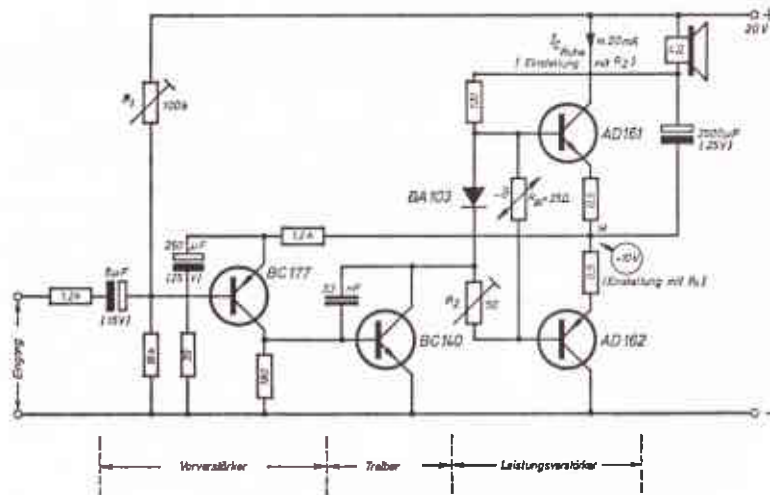


Abb. 4.56 — Schaltung des Verstärkerteils

Die Treiberstufe ist sowohl mit der Vor- als auch mit der Leistungsstufe gleichstromgegekoppelt. Da der Treibertransistor zur Ansteuerung der Leistungsstufe bereits eine verhältnismäßig hohe Leistung abgeben muß, ist er unbedingt auf ein Kühlblech oder mit Kühlstern zu montieren. Zur Vermeidung einer Schwingneigung ist zwischen Kollektor und Basis ein Kondensator (3,3 nF) geschaltet.

Auf die Treiberstufe folgt die Leistungsverstärkerstufe mit dem sehr preiswerten Transistor-Komplementärpaar AD 161/162. Damit die mit diesen Transistoren mögliche Wechselstromleistung von 10 W erreicht wird, muß die Versorgungsspannung mindestens 20 V betragen. Ebenso ist eine Kühlung der Transistoren durch Montage auf große Kühlbleche notwendig. Die Kühlbleche müssen je Transistor etwa die Maße 100 x 100 x 3 mm³ (Al) haben und können aus Platzgründen zu u-förmigen Profilen abgewinkelt werden. Die Basis-Emitter-Spannung für beide Endstufen-Transistoren wird durch eine Diode stabilisiert. Zur Temperaturstabilisierung dient ein NTC-Widerstand, der selbstverständlich direkten Wärmekontakt zum Transistorkühlblech haben muß. Die Einstellung des richtigen Arbeitspunkts für den Verstärker wird mit zwei Regelwiderständen R₁ und R₂ vorgenommen. Mit R₁ wird die sog. Mittenspannung im Punkt M eingestellt. Sie soll genau die Hälfte der Versorgungsspannung betragen. Der Regler R₂ dient zur Einstellung des Kollektor-Ruhestroms für die beiden gleichstrommäßig in Reihe geschalteten Komplementärtransistoren. Dieser Ruhestrom soll etwa 20 mA betragen und kann indirekt an einem der beiden 0,5-Ohm-Widerstände gemessen werden (0,01 V).

Die Stromversorgung für den Verstärker kann sehr einfach aufgebaut sein; denn zu den Vorzügen einer Gegentaktstufe gehört u. a. auch eine sehr geringe Brummempfindlichkeit. Es ist lediglich zu beachten,

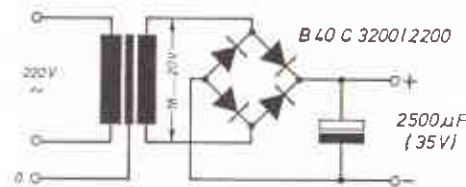


Abb. 4.57 — Stromversorgungsteil

Da die Verstärkerstufen des Regelteils äußerst empfindlich auf Brummüberlagerungen aus der Stromversorgung reagieren, ist für sie eine zusätzliche Siebung unerlässlich. Zur Siebung dienen zwei Widerstände (1,2 k Ω und 2,2 k Ω) und zwei Siebkondensatoren mit je 100 μ F Kapazität.

Hier die wichtigsten Daten des gesamten Verstärkers:

Versorgungsspannung:	20 ... 22 V
Ruhestrom ohne Signal:	ca. 50 mA
max. Betriebsstrom:	ca. 1 A
Eingangsspannung für Vollsteuerung:	ca. 50 mV
max. Ausgangsleistung:	10 W
Frequenzbereich:	30 ... 20 000 Hz
Klirrfaktor bei $P_a = 5 \text{ W}$:	$< 1 \%$
Lastwiderstand:	$\geq 4 \Omega$

Durch einen zweiten gleichartigen Verstärker ist die Erweiterung auf Stereobetrieb möglich. Dabei ist allerdings zu berücksichtigen, daß sich der Strombedarf verdoppelt. Der Netztransformator sollte für Stereobetrieb eine Belastbarkeit von ca. 3 A haben. Die in der Schaltung angegebene Gleichrichtertypen reicht auch für Stereobetrieb aus.

Die Potentiometer L, T und H werden zweckmäßig durch Stereopotentiometer ersetzt, so daß beide Kanäle gleichzeitig geregelt werden können. Für Stereobetrieb ist zudem der Einbau eines Balancereglers zu empfehlen. Er wird nach nebenstehendem Schaltungsausschnitt am einfachsten zwischen Regel- und Verstärkerteil eingesetzt. Der Balanceregler ermöglicht es, den richtigen Lautstärkeanteil aus beiden Kanälen wunschgemäß einzustellen.

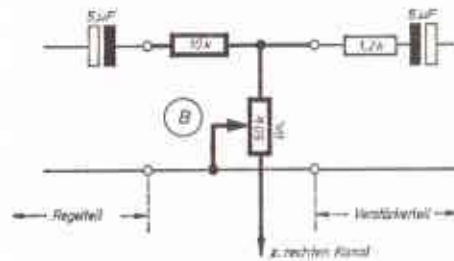


Abb. 4.58 — Einbau des Balancereglers

Für den Selbstbau empfiehlt sich die Montage der Bauteile auf gebohrten Pertinaxplatten, die auf der Unterseite mit Kupferscheiben kaschiert sind. Die fertigen Baugruppen werden zweckmäßig auf einen Metallrahmen geschraubt. Die Pertinaxplatten sollten die gleiche Länge haben.

4.5. Der Transistor in Spannungs-Stabilisierungsschaltungen

Da Referenzdioden nicht steuerbar sind und ihre Strombelastbarkeit begrenzt ist, werden zur Spannungs- und Stromstabilisierung in Stromversorgungsgeräten vielfach Transistoren eingesetzt. Ihre Wirkungsweise kann hier ganz einfach als spannungs- oder stromgesteuerter Widerstand angesehen werden.

Wir wollen die Spannungsstabilisierung mit Transistoren an einfachen Schaltungsbeispielen kennenlernen. Eine Stabilisierungsschaltung besteht im Prinzip aus einer Bezugsgröße, einem Normal, mit dem die tatsächliche Größe verglichen wird. Treten Abweichungen von der Bezugsgröße auf, so dienen diese zur Steuerung eines Reglers der den gewünschten Betriebszustand wiederherstellt. Beispielsweise dient die Zenerspannung an einer Referenzdiode als Spannungsnormal. Ist die tatsächliche Spannung größer als die Zenerspannung der Referenzdiode, so wird mit der Differenz $U - U_z$ ein Transistor als Regler gesteuert. Der Transistor begrenzt dann die Ausgangsspannung der Stabilisierungsschaltung auf den Sollwert.

4.5.1. Parallel-Stabilisierung der Ausgangsspannung

Die Abb. 4.59 zeigt eine einfache Schaltung zur Spannungs-Stabilisierung. Der Transistor liegt hier parallel zum Verbraucher und hat die Funktion eines steuerbaren Nebenschlußwiderstands.

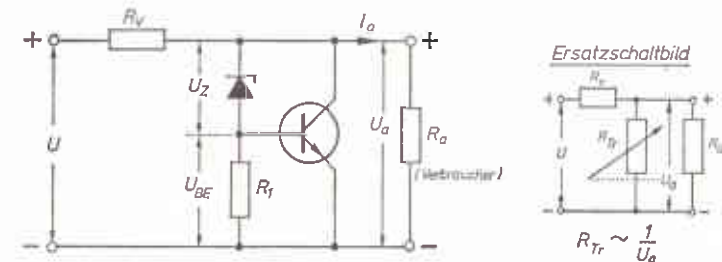


Abb. 4.59 — Parallel-Stabilisierung der Ausgangsspannung

Die Transistorstufe arbeitet wie folgt:

Der Verbraucherstrom I_a steigt. Dadurch tritt an R_v ein erhöhter Spannungsabfall auf. Die aus Referenzdiode und R_1 bestehende Spannungsteilerschaltung erhält also eine kleinere Spannung. Da aber die Zenerspannung U_z als konstant anzusehen ist, wird die Spannungsverminderung nur an R_1 wirksam. Die Basis-Emitter-Spannung des

Transistors wird damit kleiner — der Kollektorstrom sinkt. Der durch den Transistor bewirkte Nebenschluß wird hochohmiger und damit der erhöhte Spannungsabfall an R_V wieder aufgehoben.

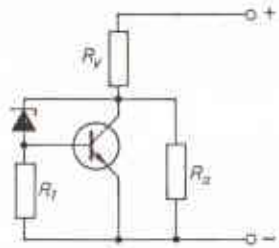


Abb. 4.60 — Darstellung als Transistorstufe mit Spannungsgegenkopplung

Die Schaltung ist als Transistorstufe mit Spannungsgegenkopplung anzusehen, wie es die nebenstehende Abb. verdeutlicht.

Der Verbraucherstrom I_a sinkt. Damit vermindert sich der Spannungsabfall an R_V . An der Spannungsteilerschaltung Referenzdiode — R_1 liegt eine höhere Spannung. Da U_z konstant bleibt, steigt die Spannung an R_1 und mit ihr die Basis-Emitter-Spannung und der Kollektorstrom des Transistors. Der Spannungsabfall an R_V wird wieder größer.

Die Versorgungsspannung U muß um den an R_V maximal auftretenden Spannungsabfall größer sein ($U \geq U_a + I_{\max} \cdot R_V$). Für einen bestimmten Änderungsbereich des Verbraucherstroms bleibt die Verbraucherspannung U_a konstant. Die Referenzdiode braucht nur mit dem Strom der Spannungsteilerschaltung belastbar zu sein. Ihre Zenerspannung bestimmt die Höhe der Ausgangsspannung. Die Ausgangsspannung beträgt $U_a = U_z + U_{BE}$. U_{BE} ist jedoch im allgemeinen sehr klein gegen U_z und kann vernachlässigt werden. Die Wirkung der Schaltung ist um so besser, je größer der Verstärkungsfaktor des Transistors ist.

Vorteil der Schaltung: Sie ist kurzschlußfest.

Nachteile: Ohne Verbraucher ist die Stromaufnahme des Transistors am höchsten. Daher ist die Schaltung nicht unbedingt leerläuft. Die Schaltung verbraucht ständig unabhängig von der Lastabnahme Energie.

4.5.2. Serien-Stabilisierung der Ausgangsspannung

Hier liegt der Regeltransistor in Serie, also in Reihe mit dem Verbraucher. Er dient als steuerbarer Vorwiderstand für den Verbraucher.

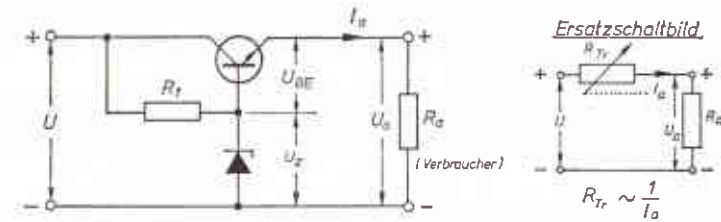


Abb. 4.61 — Serien-Stabilisierung der Ausgangsspannung

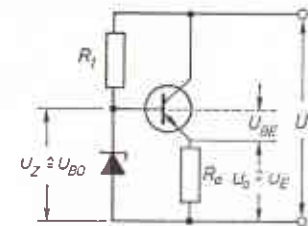


Abb. 4.62 — Darstellung als Transistorstufe mit Stromgegenkopplung

Die Schaltung hat die Wirkungsweise einer Transistorstufe mit Stromgegenkopplung; denn der Verbraucherwiderstand liegt im Emitterstromkreis.

Folgende Zustände kennzeichnen die Wirkungsweise:

Der Verbraucherwiderstand wird kleiner (I_a steigt). Der Spannungsabfall an R_a wird damit geringer. Da die Spannung U_z an der Referenzdiode konstant bleibt, steigt die Spannung $U_{BE} = U_z - U_a$. Mit steigender Basis-Emitter-Spannung U_{BE} nimmt auch der Kollektorstrom zu, der Widerstand des Transistors also ab. Der durch erhöhten Verbraucherstrom gestiegene Spannungsabfall am Transistor wird wieder auf den ursprünglichen Wert herabgesetzt.

Der Verbraucherwiderstand wird größer (I_a sinkt). Wenn R_a größer wird, steigt auch die an ihm liegende Teilspannung. Die Differenz $U_z - U_a = U_{BE}$ wird kleiner. Der Kollektor-Emitter-Strom des Transistors nimmt ab, der Widerstand des Transistors entsprechend zu. Auf die Kollektor-Emitter-Strecke des Transistors entfällt somit ein größerer Spannungsabfall. Die Verbraucherspannung bleibt konstant.

Am Ausgang der Schaltung liegt die Spannung $U_a = U_z - U_{BE}$. Da aber U_{BE} im Verhältnis zu U_z klein ist, wird die Ausgangsspannung im wesentlichen durch die Zenerspannung U_z der Referenzdiode bestimmt.

Vorteile der Schaltung: Sie ist leerlauffest. Ohne Verbraucher fließt nur ein geringer Strom über R_1 — Referenzdiode. Kein zusätzlicher Vorwiderstand erforderlich.

Nachteil: Nicht kurzschlußfest, da im Fall $R_a = 0$ am Transistor eine Basis-Emitter-Spannung in Höhe von U_z liegt und der Transistor zerstört würde.

Gef. kann die Überstromfestigkeit der Schaltung durch Einschalten eines Vorwiderstands geringfügig erhöht werden.

In der folgenden Abbildung ist die Schaltung eines kleinen spannungsstabilisierten Netzgeräts für $U_a = 9\text{ V}$ und einen Verbraucherstrom von maximal 300 mA wiedergegeben.

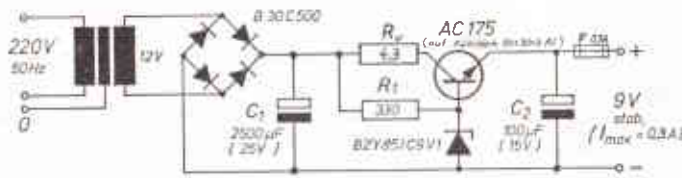


Abb. 4.63 — Schaltung eines einfachen stabilisierten Netzgeräts 9V/0,3A

Bei geringem Strombedarf kann die Kapazität C_1 kleiner gewählt werden.

Für Versuche mit Halbleiter-Bauelementen und zum Prüfen von elektronischen Schaltungen und Baugruppen ist es wünschenswert, daß die Versorgungsspannung ganz nach Bedarf eingestellt oder geregelt werden kann. Dazu ist das in folgender Schaltung dargestellte, verhältnismäßig einfach aufgebaute Netzgerät geeignet.

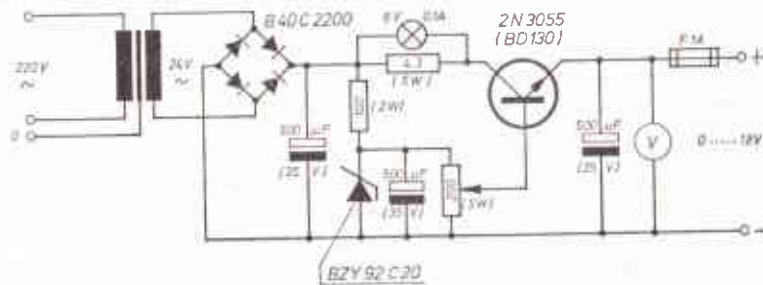


Abb. 4.64 — Schaltung eines stabilisierten Netzgeräts mit regelbarer Ausgangsspannung

Das Netzgerät arbeitet nach dem Prinzip der Serienstabilisierung. Hier ist jedoch die Spannung U_{BE} des Regeltransistors mit Hilfe eines

Potentiometers veränderbar. Die am Spannungsteilereingang liegende Spannung ist durch eine Referenzdiode stabilisiert, so daß sich Lastschwankungen nur wenig auf die Potentiometerspannung auswirken können. Als Regeltransistor dient ein preiswerter Si-Leistungstransistor 2N3055 $\approx$ BD130 ($P_v \approx 100\text{ W}$). Er muß auf ein ausreichend bemessenes Kühlblech (ca. $100 \times 100 \times 2\text{ mm}^3\text{ Al}$) montiert werden. Dem Gerät kann ein Strom von maximal 1 A entnommen werden. Die Ausgangsspannung ist beliebig zwischen fast 0 und etwa 18 V einstellbar. Eine Glühlampe dient als Überlastungsanzeige. Sie beginnt bei etwa 0,75 A Stromentnahme zu leuchten. Zur Kontrolle der Ausgangsspannung kann ein Dreheiseninstrument mit R_i ca. $2\text{ k}\Omega$ eingesetzt werden. Bei Verwendung eines Drehspulinstruments bzw. ohne Instrument ist parallel zum Ausgang ein Belastungswiderstand $1 \dots 2\text{ k}\Omega$ ($0,5\text{ W}$) zu schalten. Andernfalls läßt sich die Ausgangsspannung bei geringer Last nicht weit genug herunterregeln.

Soll das Netzgerät auch zum Betrieb von Transistorverstärkern benutzt werden, so müssen die Kapazitätswerte der Kondensatoren größer sein ($C_1 = 2500\text{ }\mu\text{F}$, C_2 und C_3 je $1000\text{ }\mu\text{F}$).

4.6. Der Transistor als Schalter

Aufgrund seiner Eigenschaften läßt sich der Transistor außer als Verstärker auch als Schalter, besser gesagt als Relais, verwenden. Zum Verständnis der Wirkungsweise eines Transistors als Schalter brauchen wir nur auf die grundsätzliche Wirkungsweise eines Transistors zurückzugreifen.

4.6.1. Die EIN-Stellung als Schalter

Liegt am Eingang eines Transistors, also zwischen Basis und Emitter, eine Spannung U_{BE} in Durchlaßrichtung, die größer ist als die Diffusionsspannung des PN-Übergangs, so fließt ein Basisstrom. Infolge der Verstärkerwirkung fließt ein erheblich stärkerer Kollektorstrom. Er ist um den Gleichstromverstärkungsfaktor β größer als der Basisstrom. Durch einen kleinen Steuerstrom I_B (z.B. 1 mA) kann ein Stromkreis mit wesentlich stärkerem Strom I_C (z.B. 100 mA) geschaltet werden. Der Transistor hat damit die Wirkungsweise eines Relais.

Damit der Transistor in EIN-Stellung als Schalter einen möglichst geringen Durchlaßwiderstand R_{CE} besitzt, muß er so stark „durchgesteuert“ werden, daß an ihm nur noch die Kollektor-Emitter-Rest-

spannung U_{CERest} bzw. U_{CEsat} abfällt (vgl. hierzu Abschn. 4.3.2). Dieser Zustand wird wieder am einfachsten anhand des Ausgangskennlinienfeldes deutlich.

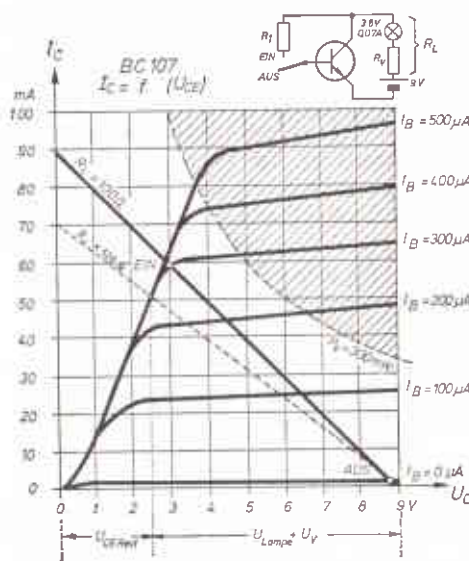


Abb. 4.65 — Der Transistor BC 107 als Schalter

viel kleineren Widerstand besitzt (Kaltwiderstand $\approx 1/10$ des Betriebswiderstands) und wählen daher einen Vorwiderstandswert von 75 Ω . Wenn der Kollektorstrom 60 mA betragen soll, muß ein Basisstrom von 300 $\mu\text{A} = 0,3 \text{ mA}$ fließen. Zu $I_B = 0,3 \text{ mA}$ gehört ein Wert von $U_{BE} \approx 0,75 \text{ V}$ (Eingangskennlinie). Somit muß der Basisvorwiderstand den Wert

$$R_1 = \frac{U - U_{BE}}{I_B} = \frac{9 \text{ V} - 0,75 \text{ V}}{0,3 \cdot 10^{-3} \text{ A}} = 27,5 \text{ k}\Omega \text{ besitzen.}$$

Im Hinblick auf ein kräftiges Durchsteuern des Transistors wird der kleinere Normwert 24 k Ω gewählt. Die Widerstandsgerade für den tatsächlichen Lastwiderstand ($R_{Lampe} + R_V$) ist im Kennlinienfeld Abb. 4.65 gestrichelt dargestellt.

In der mit den festgelegten Werten zusammengestellten Schaltung nach Abb. 4.65 leuchtet die Glühlampe, sobald der Schalter geschlossen wird. Dagegen erlischt sie beim Öffnen des Schalters.

Das Beispiel mit dem Transistor BC 107 als Schalter zeigt, daß die Restspannung $U_{CE\text{Rest}}$ im durchgeschalteten Zustand mit 2,4 V sehr

Nehmen wir an, daß mit einem Transistor BC 107 eine Glühlampe $3,8 \text{ V}/0,07 \text{ A}$ eingeschaltet werden soll. Die Betriebsspannung beträgt 9 V ($2 \times 4,5\text{-V}$ -Flachbatterie). Um eine Überlastung des Transistors sicher zu vermeiden, wird der Höchststrom auf 60 mA festgelegt. Wir zeichnen für 9 V Versorgungsspannung eine Widerstandsgerade so in das Kennlinienfeld, daß diese bei 60 mA gerade den Knick der Ausgangskennlinie schneidet. Die Widerstandsgerade trifft bei $I = 90 \text{ mA}$ auf die Stromachse. Der Lastwiderstand muß demnach

$$R_L = \frac{9 \text{ V}}{0,09 \text{ A}} = 100 \text{ } \Omega \text{ betragen.}$$

Da die Lampe im Betriebszustand einen Widerstand von

$$R_L = \frac{3,8 \text{ V}}{0,07 \text{ A}} = 54 \text{ } \Omega \text{ besitzt, ist}$$

zur Strombegrenzung zusätzlich ein Vorwiderstand von $R_V = R_L - R_{\text{Lampe}} = 100 \, \Omega - 54 \, \Omega = 46 \, \Omega$ in den Kollektorstromkreis zu schalten. Wir müssen jedoch berücksichtigen, daß die Lampe im Einschalt Augenblick einen

hoch liegt. Im angenommenen Arbeitspunkt EIN mit dem Lastwiderstand 129Ω beträgt der statische Widerstand des Transistors immerhin noch

$$R_{CE} = \frac{U_{CE\text{Rest}}}{I_C} = \frac{2,4 \text{ V}}{0,05 \text{ A}} = 48 \, \Omega !$$

Der Transistor BC 107 ist also als Schalter nicht gut geeignet. An einen Schalttransistor müssen weitaus höhere Anforderungen gestellt werden. Aus diesem Grunde weisen spezielle Schalttransistoren erheblich niedrigere Werte für die Kollektor-Emitter-Restspannung auf.

Der Betrag der Kollektor-Emitter-Restspannung läßt sich bei Schalttransistoren noch dadurch weiter herabsetzen, daß man den Transistor übersteuert, also mit einem etwas erhöhten Basisstrom durchschaltet.

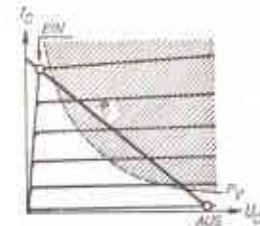


Abb. 4.68

Abb. 4.66

Bei EIN-Stellung eines Transistors als Schalter muß dieser so stark durchgesteuert werden, daß an ihm nur noch die Restspannung U_{CERest} liegt. Dazu muß ein großer Basisstrom fließen (hohe Arbeitspunkteinstellung). Je niedriger die Restspannung ist, desto kleiner ist der Widerstand des durchgeschalteten Transistors.

4.6.2. Die AUS-Stellung als Schalter

Soll der Transistor als Schalter auf AUS geschaltet werden, so ist das grundsätzlich auf drei Arten möglich:

- a) Der Basisstromkreis wird unterbrochen; es fließt also kein Basisstrom mehr. Trotzdem fließt aber noch ein, wenn auch nur geringer, Kollektorreststrom I_{CEO} , der auf die Eigenleitfähigkeit innerhalb

der PN-Übergänge B—E und B—C bei Raumtemperatur zurückzuführen ist. Aufgrund der Spannungsteilerwirkung tritt am PN-Übergang Basis-Emitter eine Teilspannung auf. Unser Transistor-Schalter ist also nicht vollständig geöffnet.

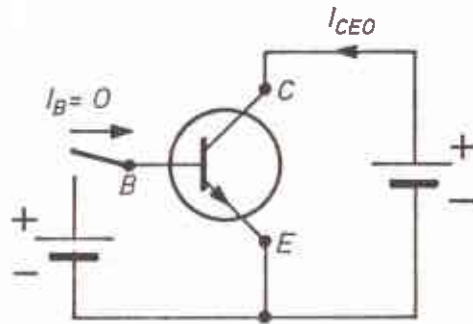


Abb. 4.67

Dieser Schaltzustand ist im Kennlinienfeld der Versuchsschaltung nach Abb. 4.65 gekennzeichnet. Für den Arbeitspunkt AUS (entsprechend $I_B = 0$) ergibt sich ein Strom von ca. 1 mA. Das bedeutet, daß der Transistor in AUS-Stellung noch einen Widerstand von

$$R_{CEAUS} = \frac{8,75 \text{ V}}{1 \cdot 10^{-3} \text{ A}} = 87,5 \text{ k}\Omega \text{ besitzt.}$$

Die Steuerung elektronischer Schalter wird in der Praxis nicht durch mechanische Schalter vorgenommen. Der mechanische Schalter ist nur zum besseren Verständnis gewählt worden. Die Steuerung eines

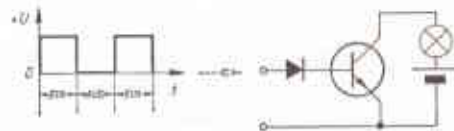


Abb. 4.68

Transistors als Schalter kann beispielsweise durch eine Rechteckspannung bewirkt werden, wie es nebenstehende Abbildung verdeutlicht. Liegt eine positive Spannung am Eingang, so ist der Transistor durchgeschaltet. Solange keine Spannung anliegt, befindet er sich in AUS-Stellung. Die Diode im Eingang dient zur Entkopplung.

- b) Der Eingang Basis-Emitter wird kurzgeschlossen; somit liegt keine Spannung am PN-Übergang Basis-Emitter. Da durch den Kurzschluß zwischen Basis und Emitter nun auch der Spannungsabfall auf dieser Strecke Null wird, ist der Kollektorreststrom I_{CES} zwar

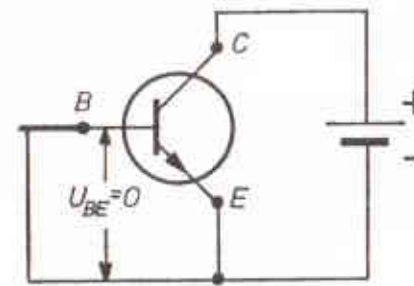


Abb. 4.69

noch kleiner als im Fall a). Ein geringer Kollektorstrom fließt aber aufgrund der Eigenleitfähigkeit immer noch. Diese Schaltung läßt sich praktisch durch eine Diode im Eingang des Transistor-Schalters verwirklichen. Zur Ansteuerung kann wieder eine Rechteckspannung mit negativen Pulsen dienen. Liegt keine Spannung am

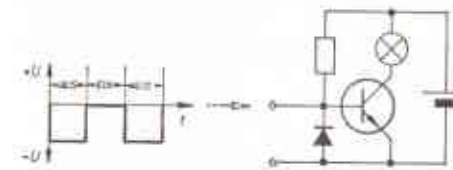


Abb. 4.70

Eingang, so wird die Diode in Sperrichtung betrieben und ist daher sehr hochohmig. Der Transistor erhält über den Basisvorwiderstand den zum Durchschalten notwendigen Basisstrom. Wird der Diode jedoch eine negative Spannung zugeführt, dann wird sie in Durchlaßrichtung betrieben. Der niedrige Durchlaßwiderstand der Diode schließt jetzt nahezu den Transistoreingang kurz. Der angestrebte

Schaltzustand wird um so besser erreicht, je niedriger die Diffusionsspannung der Diode ist; denn zwischen Basis und Emitter bleibt immer noch eine Spannung in Höhe der Diffusionsspannung der Diode wirksam. Darum werden

für solche Zwecke Germaniumdioden bevorzugt. Die Tatsache, daß ein Kurzschluß zwischen Basis und Emitter zu einem starken Absinken des Kollektorstroms führt, nutzt man in der Reparaturpraxis für Transistorschaltungen häufig aus. Soll in einer

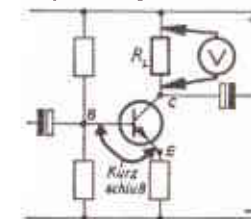


Abb. 4.71 — Funktions-

prüfung
eines Transistors
in
einer
Schaltung

fertig aufgebauten Schaltung ein Transistor auf Funktionsfähigkeit überprüft werden, so schaltet man an den Lastwiderstand R_L einen Spannungsmesser und überbrückt im Betriebszustand der Schaltung die Basis-Emitter-Strecke kurzzeitig. Bei einem funktionsfähigen Transistor sinkt dann der Kollektorstrom und somit der Spannungsabfall an R_L fast auf Null ab. Zu dieser Prüfung braucht der Transistor nicht aus der Schaltung ausgelötet zu werden! Ist der Lastwiderstand sehr klein (z.B. Spule), so kann auch die Spannung am Emitterwider-

stand gemessen werden. Da nämlich $I_E \approx I_C$, gilt für U_E ebenso: Beim Kurzschluß B—E muß U_E fast Null werden.

c) **Zwischen Basis und Emittter wird eine Spannung in Sperrichtung angelegt.** Im Gegensatz zur üblichen Schaltung der Basis-Emittter-Spannung U_{BE} in Verstärkerschaltungen wird hier der PN-Übergang Basis-Emittter in Sperrichtung betrieben. Die Sperrspannung kann so groß gemacht werden, daß kein Emittterstrom mehr fließt.

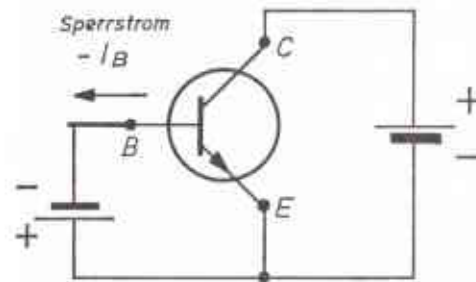


Abb. 4.72

Da wir die Eigenleitfähigkeit der Halbleiter nicht unterbinden können, fließt zwar immer noch ein kleiner Kollektorreststrom I_{CEV} , jedoch wird mit dieser Art der AUS-Schaltung der überhaupt geringstmögliche Kollektorstrom erreicht.

Die durch die drei Möglichkeiten der AUS-Schaltung erzielbaren Kollektor-Restströme stehen im Verhältnis

Unterbrechung	Kurzschluß	Sperrung
I_{CEO}	I_{CES}	I_{CEV}
wie 50	: 4	: 1.

Aus dem Zahlenverhältnis erkennt man den Grund, weshalb der Kurzschluß B—E bzw. die Sperrung durch eine Hilfsspannung in der Schaltungspraxis bevorzugt angewendet wird.

Ein praktisches Beispiel für eine Schaltungsart nach Beispiel c) ist in Abb. 4.73 wiedergegeben. Der Schalttransistor wird hier durch eine

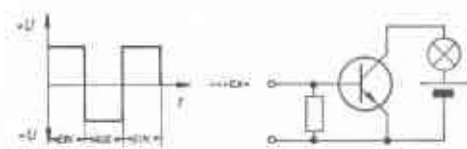


Abb. 4.73

Rechteckwechselspannung angesteuert. Dabei bewirkt positive Spannung EIN-Stellung und negative Spannung die AUS-Stellung des Transistorschalters. Der richtige Wert der Steuerspannung kann

durch Begrenzerschaltungen (z.B. mit Dioden) eingestellt werden.

In der elektronischen Datenverarbeitung wird der Transistor häufig als Nebenschlußschalter (Kurzschlußschalter) eingesetzt. Der Transistor stellt dabei den Teilwiderstand einer Potentiometerschaltung dar

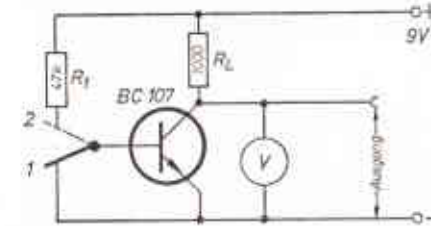


Abb. 4.74 — Meßschaltung:

Der Transistor als Schalter

und liegt parallel zum Ausgang. Die Abb. 4.74 zeigt eine einfache Meßschaltung. Mit ihr lassen sich folgende Schaltzustände erzielen:

Eingangsschalter in Stellung 1: Der Transistor ist **gesperrt** (AUS-Stellung). Sein Widerstand ist sehr groß und liegt in der Größenordnung 100 kΩ bis einige MΩ. Da nur ein sehr geringer Strom fließt, tritt am Lastwiderstand kein merklicher Spannungsabfall auf. Das Voltmeter am

Ausgang zeigt nahezu die volle Batteriespannung an.

Eingangsschalter in Stellung 2: Damit erhält der Transistor über den Basisvorwiderstand einen verhältnismäßig hohen Basisstrom. Er wird **niederohmig** (EIN-Stellung). Da jetzt ein hoher Kollektorstrom fließt, tritt am Lastwiderstand ein Spannungsabfall auf. Berücksichtigt man, daß der Widerstand des durchgeschalteten Transistors nur einige 10 Ohm beträgt, so muß der weitaus größte Teil der Batteriespannung am Lastwiderstand abfallen. Das Voltmeter zeigt daher nur noch einige hundert Millivolt an.

Diese Schaltung bewirkt eine **Umkehrung der üblichen Schalterfunktion**. Liegt nämlich am Eingang eine Spannung, so beträgt die Ausgangsspannung praktisch 0 Volt. Ist der Eingang dagegen ohne Spannung, so liegt am Ausgang die volle Betriebsspannung. In der modernen Elektronik wird diese sogenannte Umkehrschaltung am häufigsten angewendet.

4.6.3. Vergleich zwischen einem elektromechanischen Schalter und dem Transistor als Schalter

Beim Vergleich zwischen dem Transistor als Schalter und dem Schaltkontakt eines elektromechanischen Schalters (z.B. Relaiskontakt) ergeben sich zwei wesentliche Unterschiede:

- In der EIN-Stellung beträgt der Innenwiderstand unseres Transistor-Schalters **nicht 0 Ohm** wie beim mechanischen Schalter;

denn auch im durchgeschalteten Zustand haben die Halbleiterschichten einen durchaus meßbaren Widerstand, der günstigenfalls auf etwa 5...10 Ohm absinken kann. Im Ersatzschaltbild eines Transistor-Schalters müssen wir also für den Schaltzustand EIN noch seinen kleinen Durchlaßwiderstand berücksichtigen.

- b) In AUS-Stellung des Transistor-Schalters fließt in jedem Fall ein **geringer Kollektorstrom**. Unser Transistor-Schalter öffnet also den Stromkreis nie vollständig. Deshalb wird in das folgende Ersatzschaltbild eines Transistor-Schalters ein hochohmiger Widerstand parallel zu den „Schalterkontakten“ gezeichnet

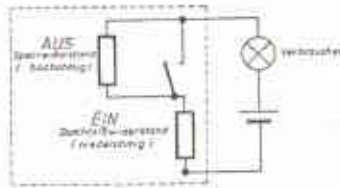


Abb. 4.75 — Ersatzschaltbild eines Transistor-Schalters

Der Transistor hat also als Schalter nicht die Eigenschaften eines idealen Schalters. In EIN-Stellung besitzt er einen nicht unbedingt vernachlässigbaren kleinen Reihenwiderstand, in AUS-Stellung einen nicht immer vernachlässigbaren hochohmigen Parallelwiderstand.

Der Vorteil des Transistor-Schalters gegenüber einem elektromechanischen Schalter liegt in seiner erheblich geringeren Baugröße und seiner um den Faktor 1000 oder mehr höheren Schaltgeschwindigkeit (keine mechanische Trägheit). Zudem treten beim Transistor-Schalter keine Verschleißerscheinungen und keine Kontaktprellungen auf.

Der Transistor findet als Schalter heute praktisch in allen elektronischen Einrichtungen breiteste Anwendung. In Schaltungen, für die der Durchgangswiderstand R_{Rest} zu hoch wird, muß er jedoch noch teilweise durch elektromechanische Schalter ersetzt werden (z.B. Durchschaltung in elektronischen Vermittlungsstellen).

4.7. Der Transistor als Schwingungserzeuger

Unter Schwingungen versteht man allgemein periodisch auftretende Zustandsänderungen. Am einfachsten sind mechanische Schwingungen vorstellbar. Wir können sie mit einer Kombination von Masse und Feder erzeugen. Hängen wir an eine Schrauben-



Abb. 4.76

feder einen Körper und stoßen diesen in senkrechter Richtung an, so schwingt er mit einer ganz bestimmten Schwingungsfrequenz auf und ab. Aufgrund der unvermeidlichen Energieverluste durch Reibung sind die erzeugten Schwingungen jedoch nur von begrenzter Dauer (gedämpfte Schwingungen). Sollen sie unvermindert weitergehen, so müssen die Energieverluste durch Energiezufuhr aufgewogen werden. Damit die Schwingungen andauern, müssen die Energieanstöße aber auch im richtigen Augenblick erfolgen. Genauer gesagt: Sie müssen gleichphasig mit der einmal erzeugten Schwingung liegen. Sonst kommen die Schwingungen sehr schnell zur Ruhe. Unser aus Schraubenfeder und Masse bestehendes Schwingungsmodell wird also nur weiterschwingen, wenn wir es jeweils im geeigneten Augenblick anstoßen. Wird der Körper beispielsweise in dem Augenblick nach unten hin angestoßen, in dem er gerade den höchsten Punkt überwunden hat, so erhält er eine zusätzliche Beschleunigung. Würde der Anstoß im gleichen Augenblick nach oben gerichtet sein, so hätte das eine Abbremsung der Bewegung zur Folge. Unser Schwingungssystem käme sehr schnell zum Stillstand.

Genauso verhält es sich mit elektrischen Schwingungen. Ein elektrisches Schwingungssystem besteht im einfachsten Fall aus einer Kapazität ($\hat{=}$ Feder bzw. Elastizität) und einer Induktivität ($\hat{=}$ Masse). Für die Reihen- oder Parallelschaltung aus Kapazität und Induktivität (= Schwingkreis) ergibt sich eine ganz bestimmte, nur von der Größe der Kapazität und der Induktivität abhängige Resonanzfrequenz. Wird ein solches elektrisches Schwingungssystem durch einen Spannungs- oder Stromstoß zu Schwingungen angeregt, so dauern diese Schwingungen nur sehr kurze Zeit (ns bis μ s). Will man die Schwingungen aufrechterhalten, so müssen die elektrischen Verluste durch ständige Energiezufuhr aufgehoben werden. Die dazu benötigte Energie kann man nun sehr einfach einem Verstärker entnehmen, an dessen Ein- oder Ausgang der Schwingkreis liegt. Ein Anteil der verstärkten Ausgangsenergie wird dabei auf den Eingang zurückgeführt. Man spricht dabei von Rückkopplung.

4.7.1. Die Rückkopplung

Die Rückkopplung kann man sich am einfachsten anhand der akustischen Rückkopplung vorstellen. Sie tritt auf, wenn sich Mikrofon und Lautsprecher einer Verstärkeranlage im selben Raum befinden.

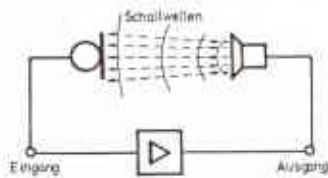


Abb. 4.77 — Akustische Rückkopplung

Die im Mikrofon in elektrische Schwingungen umgesetzten Schall-schwingungen werden verstärkt über den Lautsprecher wiedergegeben, gelangen also auch zum Mikrofon zurück, werden erneut verstärkt und so über den Lautsprecher noch lauter übertragen. Als Folge dieser „Aufschaukelung“ beginnt die Anlage zu pfeifen. Das Pfeifen dauert auch ohne weitere äußere Schalleinwirkung an. Die Frequenz des Pfeiftons hängt dabei von der Resonanzfrequenz des Lautsprechers und des

Mikrofons ab (mechanische Schwingungen). Man kann nun mit dem Lautstärkeregel eine bestimmte Einstellung finden, bei der die Schwingneigung gerade einsetzt. Bei zu geringer Lautstärkeeinstellung hört das Pfeifen vollständig auf. Wir können daraus folgern:

Damit ein rückgekoppelter Verstärker mit gleichbleibender Schwingungsstärke (Amplitude) schwingt, muß die Verstärkung genauso groß sein wie die im Rückkopplungsweg auftretenden Verluste (Dämpfung).

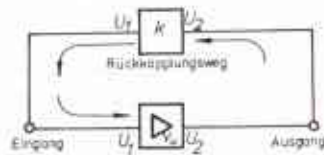


Abb. 4.78 — Blockschaltbild eines rückgekoppelten Verstärkers

Die Spannungsverstärkung beträgt

$$v_u = \frac{U_2}{U_1}$$

die „Dämpfung“ des Rückkopplungsweges (Kopplungsfaktor)

$$k = \frac{U_1}{U_2}$$

(Für den Rückkopplungsweg ist U_2 = Eingangs- und U_1 = Ausgangsspannung! Der Begriff Dämpfung ist deswegen nicht ganz korrekt, weil dem Kopplungsfaktor keine logarithmische Funktion zugrunde liegt.)

Ist die Verstärkung jetzt genauso groß wie die Verluste auf dem Rückkopplungsweg, so ergibt sich:

$$v_u \cdot k = \frac{U_2}{U_1} \cdot \frac{U_1}{U_2} = 1$$

Wir können also folgern: Damit ein rückgekoppelter Verstärker schwingt, muß die Bedingung $v_u \cdot k = 1$ erfüllt sein.

4.7.2. Mitkopplung und Gegenkopplung

Wir müssen nun noch dafür Sorge tragen, daß die Energiezufuhr über den Rückkopplungsweg

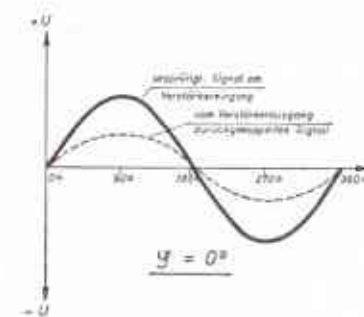


Abb. 4.79 — Mitkopplung

gleichphasig mit den einmal erzeugten Schwingungen erfolgt. **Bei einer gleichphasigen Rückkopplung spricht man von Mitkopplung.** Wie aus dem nebenstehenden Diagramm ersichtlich, haben bei der Mitkopplung die rückgekoppelte Ausgangsspannung und die Eingangsspannung gleiche Phasenlage ($\varphi = 0$). Die Spannungen addieren sich, wodurch Spannungsverluste wieder aufgehoben werden.

Wird die Rückkopplung **gegenphasig zum Eingangssignal vorgenommen, so spricht man von Gegenkopplung.** Das rückgekoppelte Signal

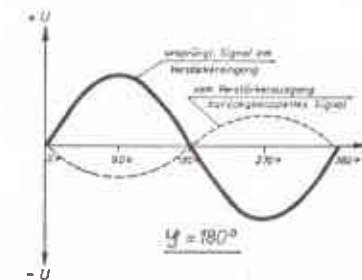


Abb. 4.80 — Gegenkopplung

wird mit einer Phasenverschiebung von 180° auf den Eingang des Verstärkers zurückgeführt. Da die Spannungen jetzt entgegengesetzt gerichtet sind, heben sie sich ganz oder teilweise auf. Die Gegenkopplung wird zur Herabsetzung von Verzerrungen in Verstärkern angewendet.

Eine einzelne Transistorstufe in Emitterschaltung hat nun für die Schwingungserzeugung den Nach-

teil, daß die Ausgangsspannung gegenüber der Eingangsspannung eine Phasenverschiebung von 180° aufweist. (Steigt die Basis-Emitter-Spannung, so sinkt die Kollektor-Emitter-Spannung

und umgekehrt.) In Transistorstufen sind daher zur Schwingungserzeugung besondere Schaltungsmaßnahmen erforderlich, um diese 180° -Phasenverschiebung wieder rückgängig zu machen.

Eine Schwingungserzeugung ist nur möglich, wenn das rückgekoppelte Signal die gleiche Phasenlage wie das Eingangssignal des Schwingungsgenerators (Schwingkreis) besitzt; Phasenbedingung $\varphi = 0$.

4.7.3. LC-Generator

Da seine Wirkungsweise am einfachsten zu verstehen ist, wollen wir uns zunächst mit einem LC-Generator beschäftigen. Die Schwingungsfrequenz wird bei ihm durch einen aus Induktivität L und Kapazität C bestehenden Schwingkreis bestimmt. Die Rückkopplung auf den Eingang kann induktiv und kapazitiv erfolgen. Beim Einschalten der Versorgungsspannung wird der Schwingkreis $L-C$ angestoßen. Auf die Sekundärspule des Übertragers wird ein Spannungsstoß induziert, der über den Kopplungskondensator C_k auf den Eingang des Transistorverstärkers zurückgekoppelt wird und verstärkt an den Schwingkreis gelangt. Es kommt zu Schwingungen mit der Resonanzfrequenz des Schwingkreises. Die Schwingungsfrequenz läßt sich mit Hilfe der Resonanzformel

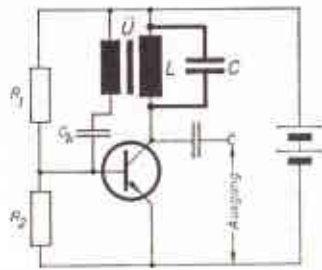


Abb. 4.81 — LC-Generator

$$f_0 = \frac{1}{2\pi\sqrt{L \cdot C}}$$

bestimmen. Am Verstärkerausgang kann die erzeugte Wechselspannung abgegriffen werden. Die in der Transistorstufe auftretende 180° -Phasenverschiebung wird durch entsprechende Polung der Wicklungsanschlüsse des Übertragers wieder aufgehoben.

Da die Induktivität einer Spule schwierig zu verändern ist, wird die Schwingungsfrequenz im allgemeinen durch eine Kapazitätsänderung beeinflusst. LC-Generatoren werden normalerweise für höhere Frequenzen gebaut (z.B. Oszillatoren in Rundfunk- und Fernsehempfängern).

Die Schaltung eines LC-Generators zur Erzeugung von Tonfrequenzen ist in Abb. 4.82 wiedergegeben. Sie enthält den uns schon bekannten

Transistor BC 107 und als Übertrager die Induktionsspule eines Fernsprechapparats W 48.

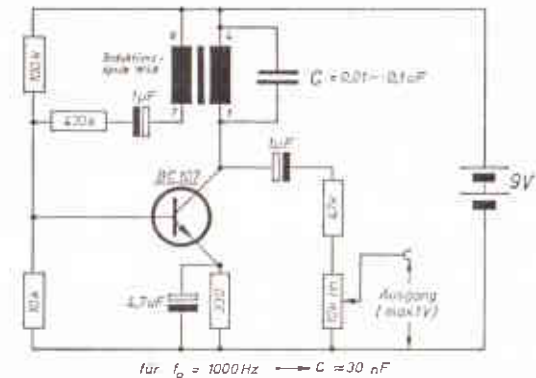


Abb. 4.82 — LC-Generator zur Erzeugung von Tonfrequenzen

Bei $C = 0,03 \mu\text{F} = 30 \text{ nF}$ beträgt die Schwingungsfrequenz etwa 1000 Hz . Da es Drehkondensatoren mit so großer Kapazität nicht gibt, ist die Frequenz nur stufenweise durch Umschalten auf andere Kapazitätswerte C veränderbar.

4.7.4. RC-Generator (Wien-Robinson-Generator/Wiengenerator)

Bei einem RC-Generator sind Widerstände (R) und Kapazitäten (C) frequenzbestimmend. Die Wirkungsweise eines RC-Generators beruht auf der Tatsache, daß er nur dann schwingt, wenn die Phasenbedingung $\varphi = 0$ zwischen Eingangs- und Rückkopplungssignal eingehalten

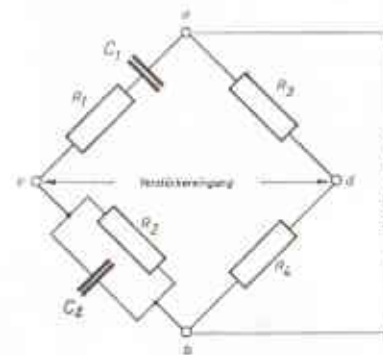


Abb. 4.83 — Wien-Robinson-Brücke

wird. Eine zur Erzeugung von Tonfrequenz-Wechselspannungen besonders gut geeignete Schaltungsart ist der Wien-Robinson-Generator. Sein Hauptmerkmal ist die Wien-Robinson-Brücke, eine Brückenschaltung aus Widerständen und Kapazitäten. Dabei liegt in einem Brückenweig (a—c) eine Reihenschaltung, in einem anderen (b—c) eine Parallelschaltung von C und R . Die Phasenbedingung $\varphi = 0$ wird beim Wien-Robinson-Generator durch zwei Schaltungs-kennzeichen bestimmt:

- Zwei hintereinandergeschaltete Transistorstufen in Emitterschaltung, von denen jede eine Phasenverschiebung von 180° verursacht, ergeben zwischen Eingangs- und Ausgangsspannung eine Phasenverschiebung von $\varphi = 360^\circ \hat{=} 0^\circ$.
- Im Brückenweig c–d tritt gegenüber der Wechselspannung im Zweig a–b **nur bei einer ganz bestimmten Frequenz keine Phasenverschiebung** auf. Die Frequenz, bei der zwischen U_{a-b} und U_{c-d} keine Phasenverschiebung vorhanden ist, lässt sich anhand der Formel

$$t_0 = \frac{1}{2 \pi \sqrt{R_1 \cdot R_2 \cdot C_1 \cdot C_2}} \quad \text{bestimmen.}$$

Diese Frequenz wird auch als Resonanzfrequenz bezeichnet; dabei sollen $R_1 = R_2$ und $C_1 = C_2$ sein.

Die Widerstände R_3 und R_4 dürfen übrigens nicht vollständig gleich sein, da sonst bei abgeglicherer Brücke die Spannung im Zweig c—d Null und somit keine Rückkopplung wirksam wäre.

Die Schwingfrequenz eines Wien-Robinson-Generators läßt sich durch gleichzeitiges Ändern der Widerstände R_1 und R_2 oder der Kapazitäten C_1 und C_2 in einem weiten Bereich ändern. Aus praktischen Gründen werden die Kapazitäten stufenweise umgeschaltet, während sich die Widerstände durch Potentiometer stufenlos regeln lassen. Wegen der Bedingung $R_1 = R_2$ werden zweckmäßig Tandem- bzw. Stereopotentiometer verwendet.

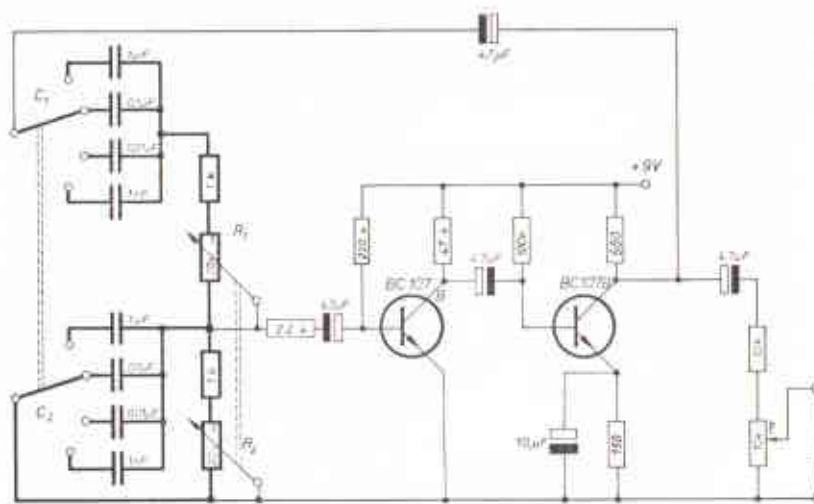


Abb. 4.84 — Wiengenerator

Im Aufbau einfacher ist der **Wiengenerator**. Als frequenzbestimmendes Glied enthält er anstelle der Wien-Robinson-Brücke einen frequenzabhängigen Spannungsteiler R_1-C_1 und $R_2 \parallel C_2$. Der zwischen Gesamtspannung und Teilspannung an dem Spannungsteiler auftretende Phasenverschiebungswinkel beträgt nur bei einer ganz bestimmten Frequenz 0° . Diese Frequenz läßt sich nach der oben angegebenen Formel berechnen. In Abb. 4.84 ist die Schaltung eines Wiengenerators für Tonfrequenzen wiedergegeben. Die verwendeten Transistoren sollten eine Gleichstromverstärkung von mindestens 250 haben, damit der Schwingungseinsatz des Generators sicher gewährleistet ist. Durch Umschaltung auf andere Kapazitätswerte und Regelung der Widerstände R_1 und R_2 lassen sich Frequenzen von etwa 20 Hz bis über den Hörbereich hinaus einstellen.

4.7.5. RC-Phasenkollengenerator

Beim Phasenkettengenerator wird durch die Reihenschaltung mehrerer RC-Glieder eine Phasendrehung von 180° erzielt. Somit ergibt sich im Rückkopplungsweg durch die 180° -Phasendrehung der Emitterschaltung und die zusätzliche 180° -Phasenverschiebung durch die RC-Kette der zum Schwingen erforderliche Verschiebungswinkel von insgesamt $360^\circ \triangleq 0^\circ$. Die Reihenschaltung eines Widerstands mit einer Kapazität ergibt eine Phasenverschiebung, die zwischen 0° und 90° liegt, jedoch auf keinen Fall 90° erreicht. Um die notwendige 180° -Phasendrehung zu erzielen, müssen also mindestens drei RC-Glieder hintereinandergeschaltet werden. Der Einfachheit halber wählt man jeweils drei gleiche Widerstands- und drei gleiche Kapazitätswerte. Die Phasenbedingung wird somit erfüllt, wenn jedes RC-Glied der Kette eine Phasendrehung von genau 60° verursacht. Aus der Formel

$$f_0 = \frac{1}{2 \pi \cdot \sqrt{6 \cdot R \cdot C}} \approx \frac{1}{15,4 \cdot R \cdot C}$$

läßt sich die Schwingfrequenz des Phasenkettengenerators bestimmen. Die in Abb. 4.85 wiedergegebene Schaltung enthält eine aus $C_1 \dots C_3$ und $R_1 \dots R_3$ bestehende Phasenkette. Die Einstellung des Arbeits-

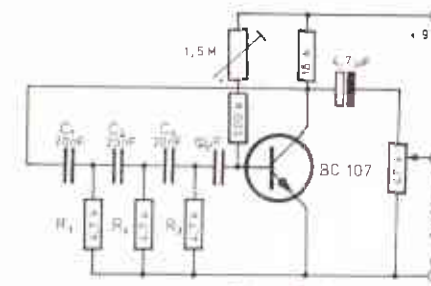


Abb. 4.85 — Phasenketten-generator

punkts für den Transistor ist dabei äußerst kritisch und wird zweckmäßig mit Hilfe eines Trimmers vorgenommen.

Der Transistor sollte eine Verstärkung von mindestens 250 besitzen. Die Schwingfrequenz des Generators beträgt

$$f_0 \approx \frac{1}{15,4 \cdot 4,7 \cdot 10^3 \Omega \cdot 20 \cdot 10^{-9} \text{ F}} = \frac{10^6}{15,4 \cdot 4,7 \cdot 20 \text{ s}} \approx 690 \text{ Hz.}$$

Der Vorteil eines einfachen Aufbaus muß beim Phasenkettengenerator mit dem Nachteil erkauft werden, daß er sich praktisch kaum in der Frequenz regeln läßt; denn dazu müßten gleichzeitig drei Kapazitäten oder drei Widerstände geändert werden. Deshalb findet man Phasenkettengeneratoren in der Praxis nur zur Erzeugung von ganz bestimmten Festfrequenzen.

4.8. Sonderformen des Transistors

Neben dem ausführlich behandelten PNP- bzw. NPN-Transistor gibt es eine Reihe von besonderen Transistorformen, die vor allem in der modernen Halbleitertechnik und elektronischen Datenverarbeitung an Bedeutung gewinnen.

4.8.1. Der Feldeffekt-Transistor (FET)

Das Prinzip des Feldeffekt-Transistors ist bereits seit etwa 1935 bekannt. Man hatte jedoch damals noch keine Kenntnis von den Vorgängen in Halbleiterstoffen oder an PN-Übergängen. So ist er erst in letzter Zeit wieder interessant geworden und wird in immer größerer Typen- und Stückzahl hergestellt. Der Feldeffekt-Transistor besteht aus einem P- oder N-dotierten Halbleiterscheibchen, das an seinen Schmalseiten zwei entgegengesetzt dotierte dünne Halbleiterschichten enthält. Wir wollen den schematischen Aufbau einmal anhand der Abb. 4.86 betrachten. Bei dem dargestellten FET handelt es sich um einen sog. N-Kanal-FET. Bei ihm besteht das Grundmaterial aus N-dotiertem Halbleiterstoff (heute meistens Silizium).

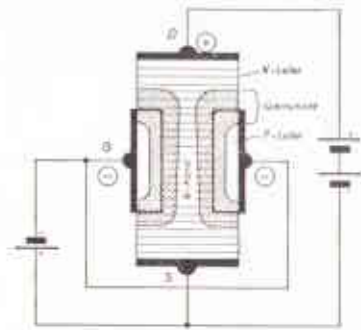


Abb. 4.86 — N-Kanal-FET

An den Seiten des FET ist je eine P-Schicht eindiffundiert worden. Die beiden P-Schichten werden leitend miteinander verbunden. Sie liegen also immer elektrisch auf gleichem Potential und stellen gemeinsam eine Anschlußelektrode des FET dar. Je eine weitere Anschlußelektrode ist am Kopfende und Fußende angebracht. Die somit vorhandenen drei Anschlüsse des FET werden wie folgt bezeichnet:

Kurzzeichen

Benennung

unterer Anschluß	S	source	(= Quelle),	
sog. Steuerelektrode	G	gate	(= Tor),	
oberer Anschluß	D	drain	(= Abfluß).	

Die Benennungen entstammen — wie nahezu alle Bezeichnungen der Halbleitertechnik — der englischen Sprache. Aus der Übersetzung erkennt man bereits die Funktion der Elektroden.

Der FET wird folgendermaßen betrieben: Zwischen S und D wird eine feste Gleichspannung angelegt. Damit fließt über den sogenannten Kanal ein bestimmter Gleichstrom. Unser N-Kanal verhält sich wie ein ohmscher Widerstand; denn in dem Stromweg S—D liegt kein PN-Übergang wie bei anderen Transistoren. Nun wird zwischen G und S ebenfalls eine Gleichspannung angelegt, und zwar so, daß der vorhandene PN-Übergang G—S bzw. G—D gesperrt ist. Legt man aber an einen PN-Übergang eine Spannung in Sperrichtung, dann entsteht an ihm eine an Ladungsträgern verarmte Zone, die sog. Sperrschicht. Die Dicke dieser Sperrschicht hängt von der Höhe der angelegten Sperrspannung ab:

große Sperrspannung	=	dicke Sperrschicht,
kleine Sperrspannung	=	dünne Sperrschicht.

Bei großer Sperrschichtdicke ist nun der leitfähige Kanal zwischen S und D erheblich enger geworden. Da der Widerstand eines Leiters von seinem Querschnitt abhängt, ist der Kanal hochohmiger geworden; es fließt ein schwächerer Strom. Wird jetzt die Sperrspannung zwischen G und S noch größer, so nimmt der Kanalstrom entsprechend ab, wird die Sperrspannung verringert, steigt der Kanalstrom. **Wir können also den Kanalstrom durch die Sperrspannung in weiten Bereichen beeinflussen.**

Da an der D-Elektrode ein höheres Plus-Potential liegt, ist übrigens die Sperrschicht in der Nähe der D-Elektrode immer etwas dicker, der leitfähige Kanal also enger.

Ein wesentlicher Vorteil des FET besteht nun darin, daß sein Steuereingang G—S in Sperrichtung betrieben wird. Er nimmt also zur Steuerung keinen Strom und somit auch keine Leistung auf. Da praktisch kein Steuerstrom fließt, ist der Steuereingang im Gegensatz zum herkömmlichen Transistor sehr hochohmig.

Bei einem Feldeffekt-Transistor kann der sogenannte Kanalstrom durch Änderung der Steuerspannung beeinflusst werden. Der FET benötigt zur Steuerung keine Leistung.

Wir haben feststellen können, daß der Feldeffekt-Transistor sowohl vom Aufbau als auch von der Funktion her viel einfacher zu beherrschen ist als der PNP- oder NPN-Transistor. Das ist auch der Grund für seine zunehmende Beliebtheit. Neben dem hier beschriebenen FET, den man wegen seines Aufbaus auch Sperrschicht-FET nennt, gibt es eine noch interessantere Form des FET, den sog. MOS-FET. Die Bezeichnung ist aus Metall-Oxid-Semikonduktor¹⁾-Feldeffekt-Transistor abgekürzt worden. Bei diesem MOS-FET ist die Steuerelektrode durch eine nichtleitende Siliziumoxidschicht isoliert angebracht. Dieses Oxid hat eine noch erheblich schlechtere Leitfähigkeit als eine PN-Sperrschicht. Der Eingangswiderstand zwischen G und S beträgt bei solchen Transistoren einige G Ω ($1\text{G}\Omega = 10^9\ \Omega$). Die Eigenschaften eines MOS-FET sind denen einer Elektronenröhre sehr ähnlich.

Die Schaltzeichen für FET sind noch nicht genormt. Um das Lesen von Schaltzeichnungen mit FET zu ermöglichen, sollen die vorläufigen Schaltzeichen dargestellt werden:

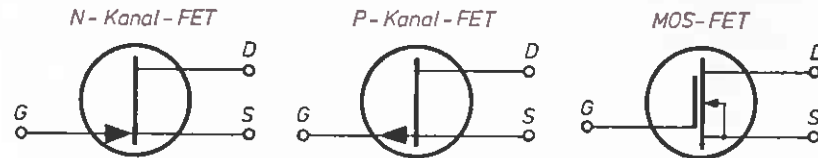


Abb. 4.87 — Schaltzeichen für FET und MOS-FET (noch nicht genormt)

In Abb. 4.88 ist die Schaltung für eine Vorverstärkerstufe mit FET dargestellt.

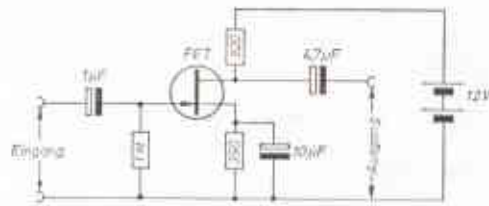


Abb. 4.88 — Einfacher Vorverstärker mit FET

4.8.2. Der Unijunktion-Transistor (UJT)

Diese auch als Doppelbasistransistor bezeichnete Art besteht aus einem schwach N-dotierten Siliziumscheibchen, an dessen Enden je

¹⁾ Semikonduktor = Halbleiter

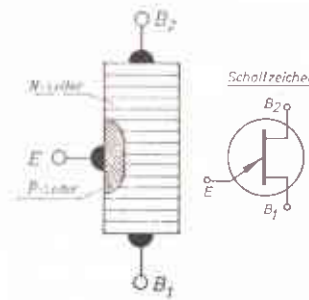


Abb. 4.89 — Schematischer

Aufbau
eines UJT

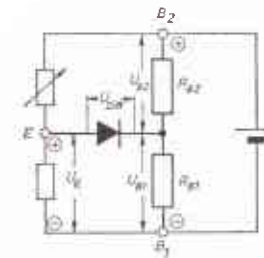


Abb. 4.90 — Ersatzschaltbild
für den UJT

ein Kontakt angebracht ist. Zwischen den äußeren Anschlüssen ist eine kleine P-leitende Zone eindotiert worden. Die beiden äußeren Anschlüsse werden als Basis 1 und Basis 2, der Anschluß an die P-Zone als Emitter bezeichnet.

Man kann sich das Ersatzschaltbild eines Unijunktion-Transistors aus einem Spannungsteiler mit zwei Teilwiderständen R_{B1} und R_{B2} sowie einer Diode $E-B_1$ zusammengesetzt denken. Wird eine Spannung zwischen B_1 und B_2 angelegt, so fließt aufgrund der schwachen Dotierung des Halbleitermaterials nur ein geringer Strom. An den beiden Teilwiderständen liegt eine Teilspannung U_{B1} bzw. U_{B2} . B_2 soll gegenüber B_1 positiv sein. Erhält der Emitter jetzt ein steigendes positives Potential, so wird der PN-Übergang $E-B_1$ niederohmig, sobald die Emitterspannung U_E größer als die Summe aus Diffusionsspannung und U_{B1} ist. Wenn sich aber der Teilwiderstand R_{B1} stark verringert hat, tritt an ihm auch nur noch eine sehr geringe Teilspannung U_{B1} auf. Die Summe $U_{Diff} + U_{B1}$ ist somit kleiner geworden. Um diesen Zustand im Transistor beizubehalten, genügt also nach dem Durchschalten auch eine viel geringere Emitterspannung U_E . Erst wenn U_E unter $U_{Diff} + U_{B1}$ abgesunken ist, wird die Emitterdiode wieder hochohmig (Talspannung).

Dieses Verhalten des Unijunktion-Transistors ermöglicht seinen Einsatz in Impulsgeneratoren, elektronischen Schaltern und Zeitgebern. Als Anwendungsbeispiel soll ein einfacher Impulsgenerator dienen (Abb. 4.91). Parallel zum Eingang $E-B_1$ des UJT liegt ein Kondensator, der über einen hochohmigen Widerstand mit dem Pluspotential der Versorgungsspannung verbunden ist. Beim Einschalten der Versorgungsspannung lädt sich der Kondensator wegen des großen Vorwiderstands nur langsam auf. Dabei wird die an ihm liegende Spannung größer.

Sobald die Kondensatorspannung ($\triangleq U_E$) die zum Durchschalten des Transistors erforderliche Höhe erreicht hat, wird die Strecke $E-B_1$ niederohmig. Über diesen niederohmigen Nebenschluß entlädt sich

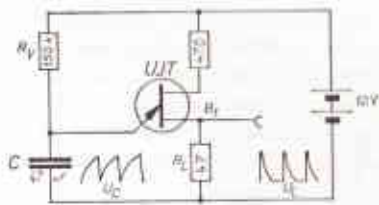


Abb. 4.91 — Impulsgenerator mit Unijunktions-Transistor

der Kondensator wieder. Wird die Talspannung unterschritten, geht der Transistor wieder in den hochohmigen Zustand über. Der beschriebene Vorgang wiederholt sich. An einem Lastwiderstand im Stromkreis E—B₁ kann die entsprechende Impulsspannung abgegriffen werden.

4.8.3. Der Fototransistor

Die Wirkungsweise aller fotoelektronischen Bauelemente beruht darauf, daß in ihnen bei Lichteinfall Ladungsträger freigesetzt werden. Wir wissen, daß in einer Fotodiode oder einem Fotoelement bei Lichteinfall eine kleine Spannung entsteht. Der Fototransistor stellt eine Kombination aus Fotodiode (bzw. Fotoelement) und Transistor dar. Der PN-Übergang Basis-Emitter dient dabei als Fotodiode. Die bei Lichteinfall entstehende Fotospannung ist als Basis-Emitter-Spannung wirksam und steuert den Transistor. Die Fotospannung ist von der Beleuchtungsstärke abhängig. Da der Kollektorstrom um den Stromverstärkungsfaktor größer ist als der Basis-Emitter-Strom, kommt hier also noch eine Verstärkerwirkung hinzu.

Der Fototransistor besitzt im einfachsten Fall nur zwei Anschlüsse, nämlich Emitter- und Kollektoranschluß; denn eine Basis-Emitter-Spannung braucht von außen nicht angelegt zu werden, da sie bei



Abb. 4.92 — Fototransistoren

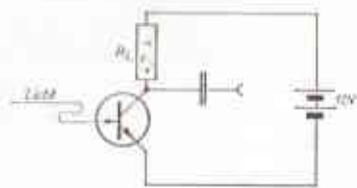


Abb. 4.93 — Schaltung eines Fototransistors

Lichteinfall im Transistor selbst erzeugt wird. Wie bei der Fotodiode ist auch in das Gehäuse eines Fototransistors eine kleine Glaslinse eingeschmolzen.

Leider ist der Arbeitspunkt eines Fototransistors stark temperaturabhängig. Um Schaltungsmaßnahmen zur Temperaturstabilisierung zu ermöglichen, besitzen einige Typen — vor allem die für höhere Leistungen — auch einen Basisanschluß (vgl. Abb. 4.92 unten). Die Grenzfrequenz der Fototransistoren liegt etwa bei 10 000 Hz, also niedriger als bei Fotodioden. Für höhere Frequenzen werden daher Schaltungen aus Fotodioden und Transistoren zusammengestellt. Der Fototransistor wird vor allem in der elektronischen Datenverarbeitung in den gleichen Anwendungsbereichen wie die Fotodiode eingesetzt.

5. Vierschichtthalbleiter-Bauelemente

5.1. Die Vierschichtdiode

Wie die Bezeichnung schon erkennen läßt, besteht die Vierschichtdiode aus vier aufeinanderfolgenden dotierten Halbleiterschichten in

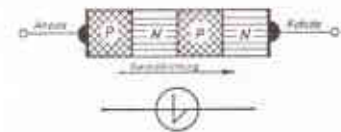


Abb. 5.1 — Schematischer Aufbau und Schaltzeichen der Vierschichtdiode

der Reihenfolge PNP, wobei nur die äußeren beiden Schichten mit einem Anschluß versehen sind. Aufgrund der Schichtfolge ergeben sich innerhalb der Vierschichtdiode drei PN-Übergänge. Legt man den Minuspol einer Spannungsquelle an die Anode, den Pluspol der Spannungsquelle an die Kathode, so verhält sich die Vierschichtdiode praktisch wie ein gesperrter PN-Übergang (also wie eine einfache Diode). Bei umgekehrter Polung der Spannungsquelle — d.h. Pluspol an der Anode und Minuspol an der Kathode — machen sich an der jetzt in Durchlaßrichtung geschalteten Vierschichtdiode gegenüber einem einfachen PN-Übergang besondere Eigenschaften bemerkbar. Betrachtet man bei der in Durchlaßrichtung geschalteten Vierschichtdiode den Zustand der drei PN-Übergänge, so ergibt sich, daß die beiden äußeren in Durchlaßrichtung liegen. Der mittlere PN-Übergang ist jedoch gesperrt. Hauptsächlich dieser mittlere PN-Übergang ist für das besondere Verhalten der Vierschichtdiode verantwortlich. Über die Vierschichtdiode beginnt nämlich erst dann ein Strom zu fließen, wenn die angelegte Spannung bis auf eine Höhe vergrößert wird, die mit der Durchbruchspannung eines einzelnen gesperrten PN-Übergangs vergleichbar ist.

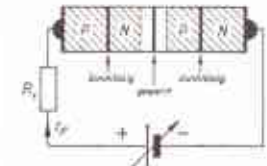


Abb. 5.2 — In Durchlaßrichtung geschaltete Vierschichtdiode

Dieser Spannungswert, bei dem es infolge Ladungsträgerdurchbruchs zu einem sehr plötzlichen Stromanstieg kommt, wird als **Zündspannung** bezeichnet. Der Durchbruch führt wie bei einer Referenzdiode nicht zur Zerstörung der Vierschichtdiode. Wesentliches Kennzeichen ihres Verhaltens ist, daß ihr Innenwiderstand nach dem Zünden schlagartig sehr klein wird. Da die Vierschichtdiode grundsätzlich mit einem Vorwiderstand betrieben werden muß, bricht die an ihr liegende Spannung nach dem Zünden in ca. 1 bis 10 μ s auf etwa 1/10 bis 1/100 der Zündspannung zusammen.

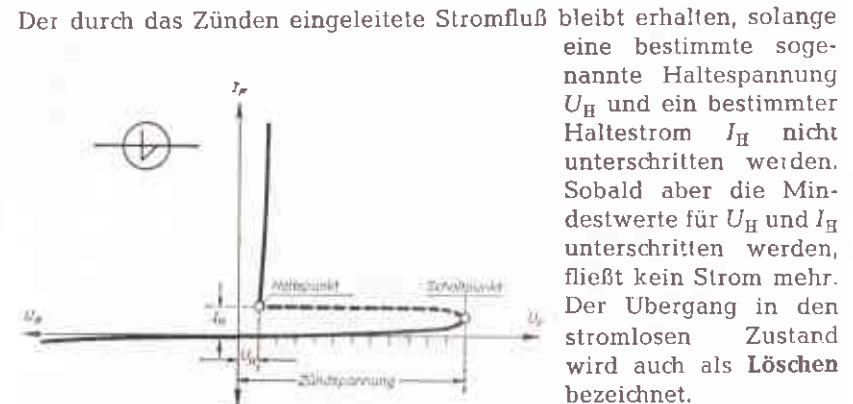


Abb. 5.3 — Kennlinie einer Vierschichtdiode

Der durch das Zünden eingeleitete Stromfluß bleibt erhalten, solange eine bestimmte sogenannte Haltespannung U_H und ein bestimmter Haltestrom I_H nicht unterschritten werden. Sobald aber die Mindestwerte für U_H und I_H unterschritten werden, fließt kein Strom mehr. Der Übergang in den stromlosen Zustand wird auch als **Löschen** bezeichnet.

Die Bezeichnungen Zünden und Löschen sind von der Glimmlampe übernommen worden, da man das Verhalten der Vierschichtdiode mit dem einer Glimmlampe ohne weiteres vergleichen kann.

Die Vierschichtdiode wird also beim Erreichen der Zündspannung sprunghaft durchlässig. Beim Unterschreiten der Haltespannung bzw. des Haltestroms wechselt ihr Verhalten ebenso plötzlich in den stromlosen Zustand (Löschen).

Wird die angelegte Spannung umgepolt, so liegen die beiden äußeren PN-Übergänge in Sperrichtung. Die Vierschichtdiode kann so praktisch nicht gezündet werden. Die Vierschichtdiode ist somit nicht gleichmäßig beeinflussbar, sondern kippt gewissermaßen in den leitenden oder nichtleitenden Zustand. Sie eignet sich daher als schneller spannungsabhängiger Schalter. Mit Hilfe von Vierschichtdioden lassen sich astabile, monostabile und bistabile Kippstufen aufbauen. Da Vierschichtdioden heute weitgehend durch Thyristoren ersetzt worden sind, wird hier auf spezielle Anwendungsbeispiele verzichtet.

5.2. Der Thyristor ¹⁾

5.2.1. Grundsätzlicher Aufbau und Wirkungsweise

Im Aufbau unterscheidet sich der Thyristor von der Vierschichtdiode im wesentlichen nur durch einen dritten Anschluß, der fast ausnahms-

¹⁾ Thyristor = Kunstwort aus Thyatron und Transistor

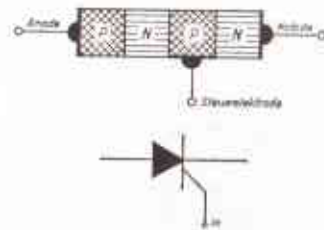


Abb. 5.4 — Schematischer Aufbau und Schaltzeichen des Thyristors

Thyristor im durchgeschalteten Zustand zu halten. Wird der Thyristor nämlich durch einen kurzen, nur einige Mikrosekunden dauernden Steuerstromimpuls gezündet, so bleibt er durchgeschaltet, solange zwischen Anode und Kathode die Haltespannung (ca. 3 V) bzw. der Haltestrom (ca. 20 ... 100 mA) nicht unterschritten wird. Um ihn wieder zu löschen — also in den stromlosen Zustand zurückzuführen — muß die Haltespannung bzw. der Haltestrom zwischen Anode und Kathode unterschritten oder bei einigen modernen Typen ein Sperrstrom über die Steuerelektrode zugeführt werden. Der Thyristor ist also eine steuerbare Vierschichtdiode. Wir sollten jedoch neben den besonderen und interessanten Eigenschaften nicht übersehen, daß der Thyristor grundsätzlich wie eine Gleichrichterdiode mit Sperr- u. Durchlaßbereich arbeitet.

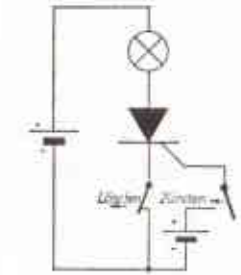


Abb. 5.5 — Prinzipschaltung eines Thyristors

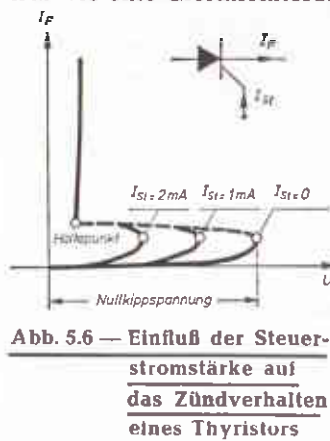


Abb. 5.6 — Einfluß der Steuerstromstärke auf das Zündverhalten eines Thyristors

los an die innere P-Zone gelegt wird. Diese dritte sogenannte Steuerelektrode ermöglicht die Beeinflussung des Zündensatzes durch erheblich geringere Spannungen bzw. Ströme als zwischen Anode und Kathode erforderlich wären. Um beispielsweise einen Thyristor mit 100 A Nennstrom und 600 V Sperrspannung zu zünden, genügt ein Strom von etwa 100 mA über die Steuerelektrode bzw. eine Zündspannung von etwa 2,5 V. Vorteilhaft ist dabei, daß der Steuerstrom nicht ständig zu fließen braucht, um den

Thyristor im durchgeschalteten Zustand zu halten. Wird der Thyristor nämlich durch einen kurzen, nur einige Mikrosekunden dauernden Steuerstromimpuls gezündet, so bleibt er durchgeschaltet, solange zwischen Anode und Kathode die Haltespannung (ca. 3 V) bzw. der Haltestrom (ca. 20 ... 100 mA) nicht unterschritten wird. Um ihn wieder zu löschen — also in den stromlosen Zustand zurückzuführen — muß die Haltespannung bzw. der Haltestrom zwischen Anode und Kathode unterschritten oder bei einigen modernen Typen ein Sperrstrom über die Steuerelektrode zugeführt werden. Der Thyristor ist also eine steuerbare Vierschichtdiode. Wir sollten jedoch neben den besonderen und interessanten Eigenschaften nicht übersehen, daß der Thyristor grundsätzlich wie eine Gleichrichterdiode mit Sperr- u. Durchlaßbereich arbeitet.

Das Zündverhalten des Thyristors geht aus Abb. 5.6 hervor. Wird kein Steuerstromimpuls zugeführt, so verhält er sich wie eine Vierschichtdiode; d.h., zum Durchschalten muß eine hohe Zündspannung (die sogenannte Nullkippspannung) angelegt werden. Je stärker aber der über die Steuerelektrode zugeführte Zündimpuls ist, desto geringer wird die zum Zünden benötigte Spannung zwischen Anode und Kathode.

Thyristoren werden mit Sperrspannungen bis zu etwa 2000 Volt und Durchlaßströmen bis zu etwa 1500 Ampere hergestellt.

5.2.2. Einfache Funktionsprüfung von Thyristoren

Um das Verhalten und die Funktionsfähigkeit eines Thyristors experimentell zu prüfen, kann man sich einer einfachen Schaltung nach

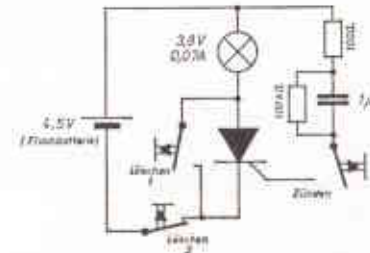


Abb. 5.7 — Versuchsschaltung zur Prüfung eines Thyristors

Abb. 5.7 bedienen. Zur Gewinnung der Zünd-Steuerimpulse wird dabei das Verhalten einer Kapazität beim Laden ausgenutzt. Wird der Schalter „Zünden“ betätigt, so lädt sich der Kondensator über den Vorwiderstand und den in Durchlaßrichtung gepolten PN-Übergang Steuerelektrode-Kathode des Thyristors auf. Da der Kondensator im Einschalt Augenblick einen Widerstand von 0 Ohm besitzt, wird der Zündstromstoß zunächst nur durch den Schutzwiderstand und den Durchlaßwiderstand des PN-Übergangs begrenzt. Der Schutzwiderstand darf nicht fehlen, da sonst der Thyristor durch einen zu hohen Steuerstrom zerstört werden könnte. Mit zunehmender Kondensatoraufladung wird der Ladestrom kleiner und sinkt nahezu auf den Wert Null. Der verhältnismäßig kurze Stromstoß reicht aus, den Thyristor zu zünden. Die Glühlampe im Thyristorstromkreis leuchtet auf und bleibt eingeschaltet, obgleich praktisch kein Steuerstrom mehr fließt. Der hochohmige Parallelwiderstand zum Kondensator dient zu dessen Entladung, damit der Zündvorgang wiederholt werden kann.

Zum Löschen wird der Thyristor zwischen Anode und Kathode durch einen Schalter „Löschen 1“ kurzgeschlossen. Infolge des Kurzschlusses liegt am Thyristor weder eine Spannung noch fließt über ihn ein Strom. Somit werden Haltespannung und Haltestrom mit Sicherheit unterschritten. Nach Aufhebung des Kurzschlusses erlischt die Lampe.

Anstatt ihn kurzzuschließen, kann man den Thyristor auch durch kurzzeitiges Unterbrechen des Stromkreises löschen (Schalter „Löschen 2“). Bei der Unterbrechung werden ebenfalls Haltespannung und Haltestrom sicher unterschritten. Nach dem Schließen des Stromkreises muß der Thyristor erneut gezündet werden.

Der mechanische Schalter zum Zünden ist in die Versuchsschaltung nur der besseren Anschauung wegen eingesetzt worden. In der Praxis wird die Zündsteuerung von Thyristoren ähnlich wie bei Schalttransistoren durch elektrische Impulse bewirkt, die in einem sogenannten Impulssteuergerät erzeugt werden und meistens nach Betrag und Phase regelbar sind. Das Löschen wird dagegen fast ausnahmslos durch Unterschreiten des Haltepunktes erreicht.

5.2.3. Anwendung von Thyristoren

Thyristoren werden überwiegend für steuerbare Gleichrichterschaltungen verwendet, z.B. auch in Stromversorgungsanlagen für Fernmeldeeinrichtungen. Daneben bieten Thyristoren die Möglichkeit einer nahezu verlustlosen Regelung der Leistungsaufnahme von Verbrauchern in Wechselstromkreisen.

Ein Thyristor kann dabei mit einem Schalter verglichen werden: Im durchgeschalteten Zustand fließt zwar der volle Strom, am Schalter liegt jedoch keine Spannung. Der Schalter selbst nimmt keine Leistung auf, weil das Produkt $U \cdot I = P = \text{Null}$ ist. Im geöffneten Zustand liegt dagegen an den Schalterkontakten die volle Spannung; aber jetzt fließt kein Strom. Die Leistungsaufnahme des Schalters, also das Produkt aus $U \cdot I$ ist auch in diesem Fall Null.

Zur Anwendung von Thyristoren zwei Beispiele:

Beispiel 1: Steuerbarer Gleichrichter:

Eine einfache Gleichrichterzelle läßt in der Einweg-Gleichrichterschaltung nur je eine Halbwelle der zugeführten Wechselstromperiode durch. Für den gewonnenen pulsierenden Gleichstrom ergibt sich ein ganz bestimmter zeitlicher Mittelwert, der bei Einweg-Gleichrichtung sinusförmigen Wechselstroms 45 % des Effektivwerts beträgt. Ein angeschlossener Verbraucher mit bestimmtem Widerstand würde eine ganz bestimmte Leistung aufnehmen. Eine hinter den Gleichrichter geschaltete Glühlampe würde beispielsweise mit einer ganz bestimmten Helligkeit leuchten.

Ein Thyristor bietet jetzt wegen seiner Steuerbarkeit die Möglichkeit, die Leistungsaufnahme des Verbrauchers zu ändern, also z.B. die Helligkeit der Glühlampe zu regeln. Zu diesem Zweck wird der Zündzeitpunkt für den Thyristor zeitlich so verschoben, daß dieser nicht mehr für die Gesamtdauer einer Halbwelle durchlässig ist. Die Wirkung dieser Maßnahme ist aus den Abb. 5.8 und 5.9 ersichtlich.

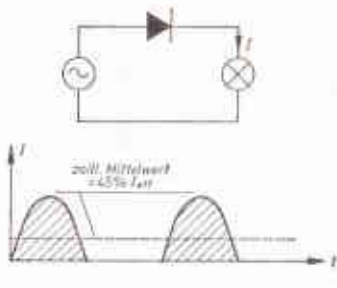


Abb. 5.8 — Pulsierender Gleichstrom nach Einweg-Gleichrichtung

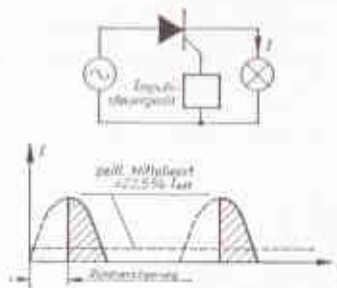


Abb. 5.9. — Pulsierender Gleichstrom nach Gleichrichtung durch einen Thyristor mit Phasenanschnittsteuerung

Abb. 5.8 zeigt das schon bekannte Strom-Zeit-Diagramm des aus einer Einweg-Gleichrichterschaltung gewonnenen pulsierenden Gleichstroms. Das in Abb. 5.9 dargestellte

Diagramm zeigt die Gleichrichtung mit einem Thyristor, wobei der Zündzeitpunkt als Beispiel so weit verzögert ist, daß der Thyristor erst bei $I_{\max}$ durchschaltet. Der Mittelwert des pulsierenden Gleichstroms ist daher auch nur halb so groß wie bei normaler Einweg-Gleichrichtung. Eine angeschlossene Glühlampe würde nur noch schwach leuchten.

Da man den Zündzeitpunkt in einem Impulssteuergerät sehr einfach in weiten Grenzen verschieben kann, ist eine Regelung zwischen voller Leistungsaufnahme und Null möglich. Dieses Verfahren, einen Thyristor erst nach einer gewissen Zeit nach Beginn der Durchlaßphase zu zünden, nennt man Phasenanschnittsteuerung. Der Vorteil dieses Regelverfahrens liegt darin, daß dabei ein hoher Wirkungsgrad erzielt wird. Nachteilig ist die durch Phasenanschnitt bewirkte Verzerrung der Sinusform. Die Verzerrung kann z.B. zu Störungen beim Rundfunkempfang führen.

Beispiel 2: Leistungsregelung im Wechselstromkreis:

Legt man zwei antiparallelgeschaltete Thyristoren in Reihe zu einem Verbraucher, so kann dessen Leistungsaufnahme auch in Wechselstromkreisen zwischen dem Höchstwert und Null geregelt werden. Die beiden Thyristoren werden dazu wieder wie im ersten Beispiel mit Phasenanschnittsteuerung durch ein Impulssteuergerät betrieben.

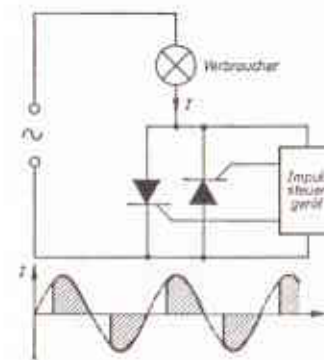


Abb. 5.10

Je ein Thyristor liegt für die positive bzw. negative Halbwelle des Wechselstroms in Durchlaßrichtung. Infolge Phasenanschnittsteuerung kommen aber die beiden Halbwellen einer Wechselstromperiode für den Verbraucher nur mehr oder weniger voll zur Wirkung, so daß die Leistungsaufnahme des Verbrauchers höher oder niedriger ist. Dieses Verfahren wird z.B. zur Regelung der Leistungsaufnahme von Beleuchtungsanlagen (Helligkeitsregelung), elektrischen Heizgeräten (Temperaturregelung) und Elektromotoren (Drehzahl- oder Leistungsregelung) angewandt.

Da phasenanschnittgesteuerte Schaltungen u.U. starke Störungen beim Rundfunkempfang verursachen, geht man heute mehr und mehr dazu über, für die Regelung von Haushaltsgeräten das Prinzip des Nullspannungsschalters anzuwenden. Dabei werden Thyristoren durch eine Regelschaltung so angesteuert, daß je nach Leistungsbedarf eine mehr oder weniger große Zahl vollständiger Wechselstromperioden gesperrt werden. Nach diesem Verfahren werden moderne Haushaltsgeräte, wie Heizgeräte, Heizkissen, Bügeleisen und Waschmaschinen, leistungsgesteuert.

5.3. Der DIAC ¹⁾ (Triggerdiode)

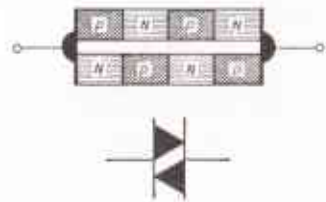


Abb. 5.11 — Schematischer Aufbau und Schaltzeichen des DIACs

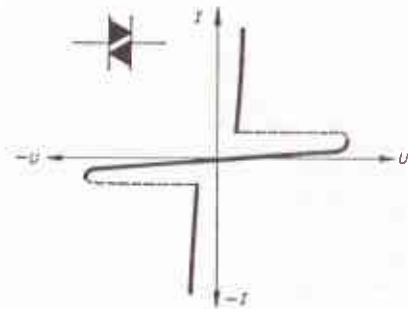


Abb. 5.12 — Kennlinie eines DIACs

lösung vom Erreichen eines bestimmten Mindestspannungswerts abhängt (vgl. hierzu auch Abschn. 1.3.1.2).

In der Fernsprechtechnik können DIACs als **Gehörschutzgleichrichter** eingesetzt werden. Da die Zündspannung eines DIACs jedoch bei 20 ... 30 V liegt, kann man sie nicht unmittelbar dem Fernhörer parallel schalten; denn 20 ... 30 V wären zur Ansteuerung viel zu viel. Trotzdem werden DIACs zum Gehörschutz eingesetzt, und zwar an der Überspannungsseite von Übertragern mit großem Übersetzungsverhältnis. Infolge des stark ausgeprägten Kippvorgangs ist die Wirksamkeit des DIACs als „Krachlöser“ erheblich höher als die des gewöhnlichen Selen-Gehörschutzgleichrichters.

¹⁾ DIAC = Kunstwort aus Diode und alternating current (Wechselstrom)

Man kann sich den DIAC vereinfacht aus zwei antiparallelschalteten Vierschichtdioden zusammengesetzt denken. Jede der beiden PNP-Schichtfolgen verhält sich wie eine Vierschichtdiode. Da sie aber in entgegengesetzter Durchlaßrichtung geschaltet sind, kann der DIAC in Wechselstromkreisen eingesetzt werden. Dem Aufbau entsprechend ergibt sich auch die Kennlinie eines DIACs aus der Zusammensetzung der Kennlinien der beiden antiparallelschalteten Vierschichtdioden. Legt man an einen DIAC eine Wechselspannung, so wird jeweils eine Vierschichtdiode des DIACs für die positive Halbwelle, die andere für die negative Halbwelle durchgeschaltet, sobald die erforderliche Zündspannung erreicht wird.

DIACs werden vorwiegend in **Triggerschaltungen** verwendet. Triggerschaltungen sind Impulsauslöser, bei denen die Aus-

Der DIAC wird in der Meßtechnik in zunehmendem Maße als **Überlastungsschutz** für Meßwerke eingesetzt. Unterhalb der Zündspannung

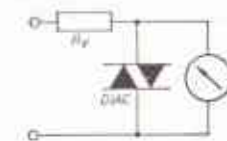


Abb. 5.13 — Überlastungsschutz mit DIAC

ist der Widerstand des DIACs so groß, daß er sich nicht verfälschend auf das Meßergebnis auswirkt. Liegt jedoch zufällig (z.B. falscher Meßbereich) eine zu hohe Spannung an, so zündet der DIAC, wird niederohmig und verhindert als Nebenschluß die Zerstörung des Meßwerks. Mit DIAC geschützte Meßinstrumente sind bis zum 100fachen Betrag des Skalenendwerts ohne Schaden für das Meßwerk überlastbar.

5.4. Der TRIAC ¹⁾

5.4.1. Grundsätzlicher Aufbau und Wirkungsweise

Zur Ausnutzung der Thyristoreigenschaften in Wechselstromkreisen — z.B. zur Leistungsbegrenzung — sind zwei antiparallelschaltete Thyristoren erforderlich; denn ein Thyristor liegt nur jeweils für eine Halbwelle des Wechselstroms in Durchlaßrichtung. Nachteilig ist dabei, daß jeder der beiden Thyristoren für seine Durchlaßphase mit einem Steuerimpuls durchgeschaltet werden muß. Dazu ist ein Impulssteuergerät mit zwei Ausgängen erforderlich. Dieser Nachteil

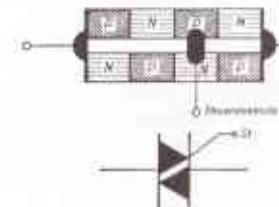


Abb. 5.14 — Schematischer Aufbau und Schaltzeichen des TRIACs

wird vermieden, wenn man **zwei antiparallelschaltete Thyristoren zu einem Bauelement zusammenfaßt und mit nur einer Steuerelektrode** versieht. Ein so aufgebautes Bauelement wird als **TRIAC** bezeichnet. Die Steuerelektrode liegt dabei für einen der beiden Thyristoren an der inneren P-Zone, für den anderen an der inneren N-Zone. Wird nun über die Steuerelektrode ein Impuls bestimmter Stromrichtung zugeführt, so wird jeweils eine der Vierschichtstrecken durchgeschaltet und die andere gesperrt. Bei entgegengesetzter Richtung des Steuerimpulses kippen beide Vierschichtstrecken ebenso in den entgegengesetzten Schaltzustand.

Die Kennlinie des TRIACs ergibt sich aus der Zusammensetzung der

Die Kennlinie des TRIACs ergibt sich aus der Zusammensetzung der

¹⁾ TRIAC = Kunstwort aus Triode und alternating current (Wechselstrom)

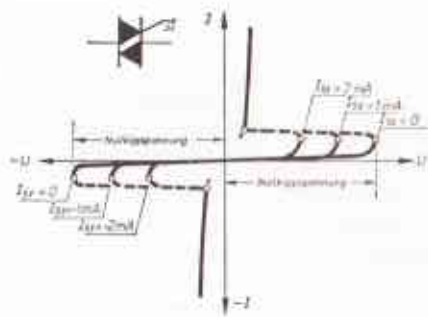


Abb. 5.15 — Kennlinie eines TRIACs

Kennlinien beider antiparallelschalteter Thyristoren. Da der TRIAC für den Einsatz in Wechselstromkreisen vorgesehen ist, entfallen besondere Schaltungsmaßnahmen zum Löschen. Die jeweils gerade durchlässig geschaltete Strecke wird automatisch dadurch gelöscht, daß in der Nähe des Nulldurchgangs der Wechselspannung bzw. des Wechselstroms mit Sicherheit Haltespannung und Haltestrom unterschritten werden.

5.4.2. Anwendung des TRIACs

Der TRIAC wird fast ausnahmslos zur Leistungsregelung in Wechselstromkreisen eingesetzt. Wegen seiner einfachen Steuerbarkeit ist nur ein geringer Schaltungsaufwand zur Impulssteuerung erforderlich. Die Impulssteuerung kann z.B. durch Kippschaltungen mit Glühlampen oder Unijunktion-Transistoren oder einfach durch einen DIAC bewirkt werden.

Dazu ein Beispiel:

Leistungsregelung durch Phasenanschnittsteuerung: Wir wollen dazu zunächst eine einfache Grundschialtung (Abb. 5.16) betrachten. In der Regelschaltung liegt der TRIAC in Reihe mit dem Verbraucher (Lampe). Die Nullkippspannung des TRIACs (d.h. die zum Zünden erforderliche Spannung, wenn kein Steuerimpuls zugeführt wird) muß so hoch liegen, daß er ohne Steuerimpuls nicht etwa durch die Spannungsspitzen der

Wechselspannung gezündet werden kann. Das bedeutet für die Anwendung in 220-Volt-Netzen, daß die Nullkippspannung über 311 V liegen muß. Praktisch wählt man dafür TRIACs mit etwa 400 V Nullkippspannung. Zur Steuerung nutzt man das Verhalten einer Kapazität und eines DIACs aus. Steigt die angelegte Spannung, so lädt sich der Kondensator auf. Der zeitliche Anstieg der Kondensatorspannung hängt dabei von der Größe des Vorwiderstands ($R_V + R_I$) und der Kapazität des Kondensators ab (vgl. hierzu Abschn. 7). Der DIAC kann erst einen Steuerimpuls für den TRIAC liefern, wenn die Ladespannung des Kondensators ihren Zündspannungswert von etwa 30 Volt erreicht hat.

Der Vorgang wiederholt sich für die entgegengesetzte Richtung der angelegten Spannung.

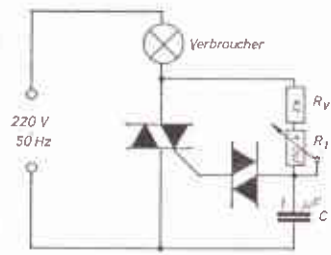


Abb. 5.16 — Grundschialtung

Mit Hilfe des Regelwiderstands R_I läßt sich der zeitliche Anstieg der Kondensatorspannung (Zeitkonstante) beeinflussen. Je größer R_I ist, desto langsamer lädt sich der Kondensator auf und um so später wird auch der Zündspannungswert des DIACs erreicht. Der Schutzwiderstand R_V soll den Steuerstrom begrenzen, so daß DIAC und TRIAC nicht überlastet werden können.

Die beschriebene Grundschialtung hat jedoch den Nachteil, daß mit ihr keine gleichmäßige Abwärtsregelung der Leistungsaufnahme bis Null möglich ist. Wird der Widerstand R_I immer weiter vergrößert, so geht die Aufladung des Kondensators schließlich zu langsam vor sich. Bevor der Zündspannungswert des DIACs erreicht wird, ist die angelegte Wechselspannung schon wieder abgesunken. Bei einer bestimmten Stellung des Reglers R_I wird der TRIAC also nicht mehr durchgesteuert und die Leistungsaufnahme des Verbrauchers fällt sprunghaft auf Null. Dieser Nachteil wird durch die erweiterte Schaltung nach Abb. 5.17 vermieden. Auf die Erläuterung funktionsmäßiger Einzelheiten muß in diesem Rahmen verzichtet werden. Grundsätzlich besitzt die Schaltung die gleiche Wirkungsweise wie die Grundschialtung.

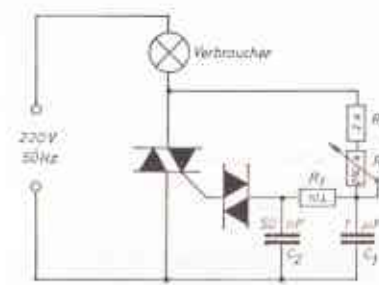


Abb. 5.17 — Verbesserte Regelschaltung

TRIAC und DIAC werden heute schon vielfach zusammen in ein Gehäuse eingebaut.

Zusammenfassend kann festgestellt werden: **Wesentliches Merkmal aller Vierschichtbleiter-Bauelemente ist das sprunghafte Kippen in den leitenden oder nichtleitenden Zustand.** Der Kippvorgang kann dabei entweder durch die an die äußeren Elektroden angelegte Spannung (Vierschichtdiode und DIAC) oder durch niedrige Spannungs- oder Stromwerte über eine besondere Steuerelektrode (Thyristor und TRIAC) eingeleitet werden. Durch Antiparallelschaltung sind die Bauelemente auch für Anwendungen in Wechselstromkreisen geeignet.

Vierschichtbleiter-Bauelemente werden für Kippstufen aller Art, gesteuerte Gleichrichterschaltungen und Leistungsregelschaltungen verwendet. Mit ihnen lassen sich hohe Wirkungsgrade erzielen, da sie selbst nur wenig Leistung aufnehmen.

6. Elektronenröhren

Im Zuge der stürmischen Entwicklung der Elektronik wurde die Elektronenröhre durch Halbleiter-Bauelemente immer weiter verdrängt. Trotzdem gibt es Anwendungsbereiche, in denen die Elektronenröhre bis heute noch nicht durch gleichwertige oder bessere Halbleiter-Bauelemente ersetzt werden konnte. Dazu gehören z.B. die Leistungsstufen in Sendern für Rundfunk-, Fernseh- und Fernspreübertragung. Mit Elektronenröhren lassen sich Leistungen bis zu einigen Hundert Kilowatt erreichen, während die Leistungsgrenze für Transistoren immer noch bei einigen Hundert Watt liegt. Zudem liegt die mögliche Grenzfrequenz bei Elektronenröhren höher als bei Halbleiter-Bauelementen, so daß auch in der Hochfrequenz- bzw. Höchsthäufigkeitstechnik Elektronenröhren noch eine erhebliche Rolle spielen.

6.1. Grundsätzlicher Aufbau und Wirkungsweise der Elektronenröhre

Die Elektronenröhre — kurz Röhre genannt — besteht im allgemeinen aus einem luftleergepumpten Glaskolben, in den Elektroden eingeschmolzen sind. Da das Vakuum elektrisch nicht leitfähig ist, muß die Leitfähigkeit des Röhreninneren durch Einbringen von Ladungsträgern erzielt werden. Das geschieht durch ein elektrisches Heizsystem, kurz Heizung genannt. Die Wirkungsweise der Heizung beruht auf der Tatsache, daß ein stark erwärmter Leiter Elektronen aus seiner Oberfläche herausschleudert (emittiert). Den Vorgang bezeichnet man als Emission. Material des Heizsystems und Temperatur beeinflussen die Stärke der Emission. Da aber das besonders temperaturbeständige Heizleitermaterial Wolfram keine besonders gute Emissionseigenschaften besitzt, bringt man unmittelbar oder auf einen besonderen Träger über dem Wolfram-Heizfaden ein emissionsfreudiges Material auf (meistens Bariumoxid oder Strontiumoxid). Je nach Aufbau des Heizsystems unterscheidet man zwischen der direkten und der indirekten Heizung.

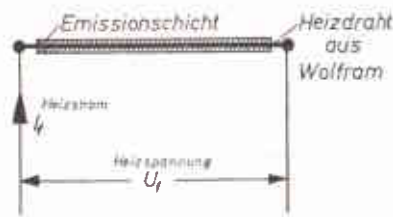


Abb. 6.1. — Direkte Heizung

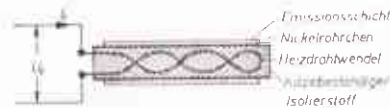


Abb. 6.2 — Indirekte Heizung

Bei der **direkten** Heizung liegt unmittelbar auf dem Wolfram-Heizfaden eine Emissionsschicht aus Barium- oder Strontiumoxid. Wegen des unmittelbaren Wärmeübergangs vom Heizfaden auf die Emissionsschicht ist der Heizleistungsbedarf der direkten Heizung verhältnismäßig gering. Damit die Elektronenemission nicht schwankt, muß die direkte Heizung mit reinem Gleichstrom betrieben werden. Die Anheizzeit beträgt weniger als eine Sekunde.

Kennzeichen der **indirekten** Heizung ist die galvanische Trennung von Emissionsschicht und Heizfaden. Die Emissionsschicht liegt hier auf einem Nickelröhrchen, das vom eigentlichen Heizfaden isoliert ist. Da elektrische Nichtleiter meistens auch schlechte Wärmeleiter sind, bedingt der Betrieb einer indirekten Heizung einen verhältnismäßig hohen Leistungsbedarf. Wegen der galvanischen Trennung kann jedoch der Heizfaden ohne wesentliche Beeinflussung der Röhrenfunktion mit Wechselstrom betrieben werden. Dem kommt auch die große Wärmeläufigkeit des Heizsystems entgegen; denn die Anheizzeit der indirekten Heizung beträgt etwa 30 bis 40 Sekunden. Damit wirken sich selbst langsame Schwankungen des Heizstroms kaum auf die Elektronenemission aus.

Bringt man in einem luftleergepumpten Glaskolben gegenüber der Heizung eine Elektrode an, so wird bei Erwärmung des Heizfadens ein Teil der emittierten Elektronen zu der Gegenelektrode fliegen. Ihre Zahl ist jedoch äußerst gering, da die negativ geladene Elektronenwolke (Raumladungswolke) über dem Heizsystem den Austritt weiterer Elektronen erschwert und die Geschwindigkeit der Elektronen noch verhältnismäßig gering ist. Bringt man jedoch in den äußeren Stromkreis eine Spannung, das Pluspotential an der Gegenelektrode, so werden fast alle emittierten Elektronen von der positiven Elektrode (Anode) angezogen.

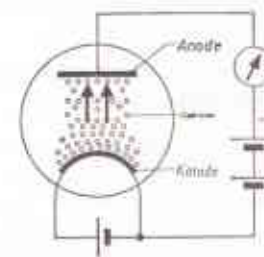


Abb. 6.3 — Elektronenbewegung
in einer Röhre

Die Bewegungsgeschwindigkeit der Elektronen hängt von der Spannung zwischen Anode und Kathode ab:

$$v = 600 \sqrt{U} \frac{\text{km}}{\text{s}}, \text{ d.h. sie beträgt}$$

$$\text{bei 1 V ca. } 600 \frac{\text{km}}{\text{s}},$$

$$\text{bei 100 V ca. } 6000 \frac{\text{km}}{\text{s}},$$

$$\text{bei 10 000 V ca. } 60 000 \frac{\text{km}}{\text{s}}.$$

Keht man die Polarität der angelegten Spannungsquelle um, so können keine Elektronen vom Heizfaden zur anderen Elektrode gelangen, da sich gleiche Ladungen abstoßen. Die Elektronenröhre läßt den Elektronenstrom also nur in einer Richtung durch.

6.1.1. Zweipolröhre (Diode)

Eine Röhre mit zwei Elektroden, dem Heizsystem (Katode) und der Gegenelektrode (Anode), bezeichnet man als **Zweipolröhre** (Diode). Trägt man den Röhrenstrom in Abhängigkeit von der angelegten Spannung in ein Diagramm, so erkennt man an der Kennlinie:

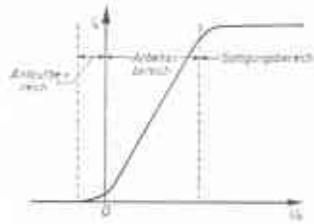


Abb. 6.4 — Kennlinie einer Zweipolröhre

a) Schon bei einer Anodenspannung von 0 V fließt ein geringer Strom durch die Röhre. Er hat eine Stärke von ca 1 mA und kommt durch diejenigen Elektronen zustande, die bei der Elektronenemission stark genug beschleunigt worden sind.

Man nennt diesen Strom **Anlaufstrom**. Er wird in der modernen Röhrentechnik dazu benutzt, in Anfangsstufen die Gittervorspannung einer Röhre zu erzeugen.

b) Wird an die Anode eine positive Spannung gelegt, so steigt mit wachsender Spannung der Anodenstrom der

Röhre. Innerhalb eines bestimmten Bereichs ist der Anstieg der Kennlinie nahezu geradlinig. Diesen Bereich der Kennlinie bezeichnet man als **Arbeitsbereich**.

c) Steigert man die Höhe der Anodenspannung immer weiter, so ist schließlich keine Steigerung des Anodenstroms mehr zu erreichen. Grund: Sämtliche von der Katode emittierten Elektronen werden sofort von der Anode aufgenommen. Diesen Kennlinienbereich nennt man **Sättigungsbereich**.

Bei modernen Elektronenröhren kann der Sättigungsbereich nicht erreicht werden, da er wegen der hochentwickelten Emissionsschichten so hoch liegt, daß die Röhre zerstört würde.

d) Kehrt man die Polarität der Anodenspannung um, so fließt kein Anodenstrom. Grund: Die negativ geladene Anode stößt die ebenfalls negativ geladenen Elektronen ab.

Wegen der Eigenschaft, den Elektronenstrom nur in einer Richtung durchzulassen, wird die Zweipolröhre als Gleichrichter verwendet.

Voraussetzung für das ungestörte Fließen des sogenannten Anodenstroms ist ein ausreichendes Vakuum in der Röhre. Etwa noch vorhandene Gasreste werden ionisiert. Dadurch kommt es zu einer Gasentladung (Glimmentladung), die zur Zerstörung der Röhre führen kann. Zur chemischen Bindung der Gasreste wird in die Röhren ein sogenannter „**Getter**“ eingebaut. Wird die Röhre zu stark belastet, so können außerdem durch die Erwärmung des Anodenblechs Gase frei werden, die das einwandfreie Arbeiten der Röhre beeinträchtigen.

6.1.2. Dreipolröhre (Triode)

Mit Hilfe einer dritten Elektrode, die zwischen Heizsystem (Katode) und Gegenelektrode (Anode) angebracht wird, kann man den Elek-

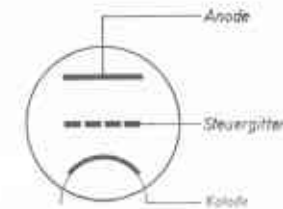


Abb. 6.5 — Dreipolröhre (Triode)

tronenstrom in der Röhre beeinflussen (steuern). Diese meistens wendelförmige Drahtelektrode wird als **Steuergitter** bezeichnet.

Da sich das Steuergitter (auch kurz „Gitter“ genannt) in unmittelbarer Nähe der Katode befindet, beeinflusst es mit seiner Spannung den Elektronenstrom.

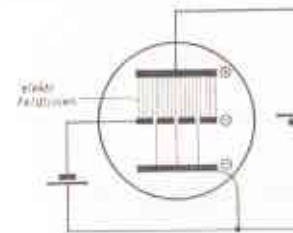


Abb. 6.6 — Durchgriff

Die Anodenspannung kann dagegen nur mit einem Bruchteil ihres Betrages wirken, weil nur ein Teil der elektrischen Feldlinien durch das Gitter „hindurchgreift“.

Den für die Stärke des Anodenstroms wirksamen Anteil der Anodenspannung bezeichnet man daher auch als „**Durchgriff**“ **D**. Er wird in % der Anodenspannung angegeben.

Für die Stärke des Elektronenstroms in der Röhre sind also maßgebend:

- die Spannung am Steuergitter und
- der wirksame Anteil der Anodenspannung.

Legen wir an das Steuergitter der Röhre eine genügend hohe negative Spannung, so werden die von der Katode emittierten negativen Elektronen durch das Gitter abgestoßen. Es fließt somit kein Strom durch die Röhre. Läßt man jetzt die Gitterspannung langsam Null werden, dann wird der Elektronenstrom immer stärker zur Anode fließen, da die hemmende Wirkung des Gitters nachläßt.

(Achtung! Das Gitter darf im allgemeinen nicht positives Potential haben, da dann der Anodenstrom so stark wird, daß die Röhre zerstört wird. Daher muß das Gitter immer eine negative Vorspannung erhalten.)

Die Stärke des Elektronenstroms hängt von der elektrischen Feldstärke in der Umgebung der Katode ab. Die Steuerwirkung des Gitters wird also um so größer, je geringer der Abstand Gitter-Katode gemacht wird. In Röhren mit sog. Spannungsgitter wird ein äußerst geringer Gitter-Katode-Abstand von etwa $30\text{ }\mu\text{m}$ und damit eine große Verstärkerwirkung erreicht.

Um die Steuerwirkung des Gitters zu erkennen, stellt man meistens für eine Dreipolröhre die sogenannte I_a/U_g -Kennlinie dar. Dazu wird der Anodenstrom I_a in Abhängigkeit von der Spannung U_g gemessen. Die Anodenspannung bleibt dabei konstant.

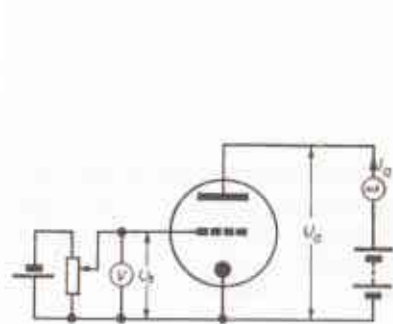


Abb. 6.7 — Meßschaltung zur Aufnahme der I_a / U_g -Kennlinie

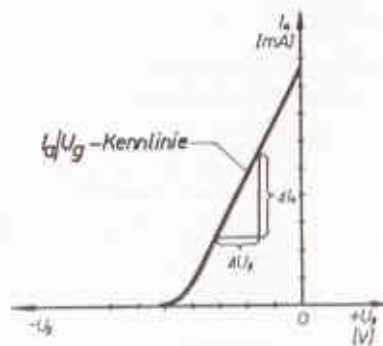


Abb. 6.8 — I_a / U_g -Kennlinie

Da in der Röhre keine Elektronen zum negativen Steuergitter gelangen (gleiche Ladungen stoßen sich ab), kann im Gitterkreis auch kein Strom fließen.

Daraus folgt:

Zur Steuerung einer Verstärkerröhre wird im Gitterkreis keine Leistung verbraucht.

Je steiler die I_a/U_g -Kennlinie verläuft, um so stärker ändert sich der Anodenstrom I_a bei gleicher Veränderung der Gitterspannung U_g .

Man bezeichnet das Verhältnis von $\frac{\Delta I_a}{\Delta U_g}$ als „Steilheit“ S

$$S = \frac{\Delta I_a}{\Delta U_g} \quad \frac{\text{mA}}{\text{V}}$$

einer Röhre (vgl. hierzu Abb. 6.8).

Die Steilheit drückt aus, um wieviel mA sich der Anodenstrom ändert, wenn die Gitterspannung um 1 V verändert wird. Ihre Maßeinheit ist entsprechend mA/V. Für die Praxis interessiert dabei nur die Steilheit im geradlinigen Verlauf der Kennlinie (sogenannter linearer Bereich). Die Steilheit ist ein wichtiges Merkmal für die Arbeitsweise der Röhre. Bei modernen Röhren erreicht man Röhrensteilheiten bis zu etwa 40 mA/V.

6.2. Elektronenröhre als Verstärker

Um die Verstärkerwirkung einer Elektronenröhre ausnutzen zu können, muß sie an ihrem Gitter von außen gesteuert und in den Röhrenstromkreis ein Verbraucher R_a geschaltet werden. Die schematische Schaltung einer Röhrentriode als Verstärker ist in Abb. 6.9 dargestellt.

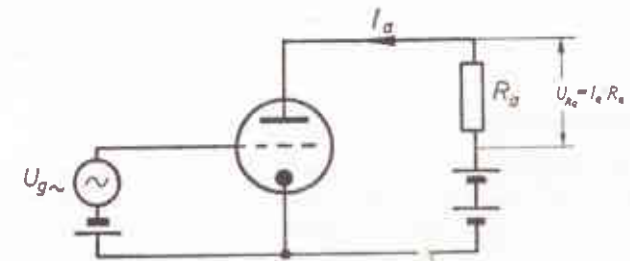


Abb. 6.9 — Triode als Verstärker

Da das Steuergitter nicht positives Potential erhalten darf, muß eine feste negative Gittervorspannung angelegt werden. Die Gittervorspannung legt den Arbeitspunkt auf der I_a/U_g -Kennlinie fest und ermöglicht es, die Röhre durch eine Wechselspannung anzusteuern.

Ohne Steuerung fließt über die Röhre und den Verbraucherwiderstand R_a ein Ruhestrom I_a , der an R_a einen Spannungsabfall $U_a = I_a \cdot R_a$ hervorruft. Wird an das Steuergitter zusätzlich zur negativen Vorspannung eine Wechselspannung gelegt, so steigt der Anodenstrom mit der positiven und fällt mit der negativen Halbwelle der Steuerwechselspannung. Damit ergeben sich Änderungen des Spannungsabfalls an R_a . Diese Spannungsänderungen $\Delta U_{R_a} = \Delta I_a \cdot R_a$ sind im allgemeinen wesentlich größer als die Änderungen der Gitterspannung. Die Triode bewirkt somit eine Spannungsverstärkung, die

als **Verstärkungsfaktor** angegeben wird und sich nach folgender Formel berechnen läßt:

$$\nu = \frac{\Delta U_{Ra}}{\Delta U_g}.$$

Der Verstärkungsfaktor sagt aus, um wieviel Volt sich die Spannung an R_a ändert, wenn die Gitterspannung U_g um 1 Volt verändert wird. Er ist von der Größe des Verbraucherwiderstands abhängig und wird daher in Tabellenbüchern immer für ganz bestimmte Widerstandswerte von R_a und einen bestimmten Arbeitspunkt angegeben. Die Verstärkungsfaktoren moderner Trioden liegen etwa zwischen 20 und 100. Das Besondere am Röhrenverstärker ist, daß er im Gegensatz zum Transistorverstärker leistungslos — nur durch eine Spannung — gesteuert werden kann.

Jede Röhre hat aufgrund ihrer Konstruktionsmerkmale einen bestimmten inneren Widerstand R_i . Zur Erzielung der **größten Verstärkerleistung** gilt auch bei der Wahl des Widerstandswertes für R_a :

$$R_a = R_i \text{ Röhre}.$$

In der Praxis wird R_a jedoch im allgemeinen zur Herabsetzung der Verzerrungen größer gewählt als R_i der Röhre.

6.2.1. Arbeitspunkteinstellung bei Röhrenverstärkern

Da das Steuergitter-Potential einer Elektronenröhre negativ, das Anodenpotential dagegen positiv sein muß, kann die Gittervorspannung nicht einfach der Anodenspannungsquelle entnommen werden. Zur Erzeugung der Gittervorspannung — eigentlich zur Arbeitspunkteinstellung — sind drei schaltungstechnische Maßnahmen üblich:

a) Gewinnung der Gittervorspannung aus einer besonderen Gitterspannungsquelle

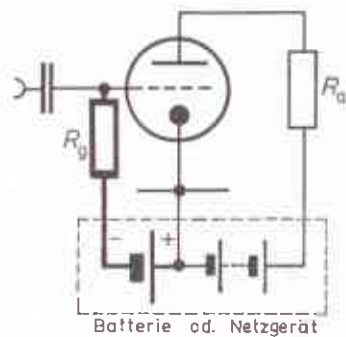


Abb. 6.10

Bei netzgespeisten Röhrenverstärkern läßt sich die negative Gittervorspannung aus dem Netzgerät entnehmen, wobei allerdings sehr hohe Anforderungen an die Siebung gestellt werden; denn die Röhre verstärkt jede Spannungsschwankung. In batteriebetriebenen Geräten werden eine bis mehrere Zellen der Anodenbatterie zur Erzeugung der negativen Gittervorspannung ausgenutzt.

Vorteile: Feste, belastungsunabhängige Arbeitspunkteinstellung. Verbindung Katode-Masse möglich (geringe Brummneigung).

Nachteile: Erhöhter Schaltungsaufwand bzw. erhöhter Zellenbedarf bei Batteriebetrieb. Keine Schutzwirkung bei Überlastung.

b) Einschalten eines Widerstands in die Katoden-zuleitung

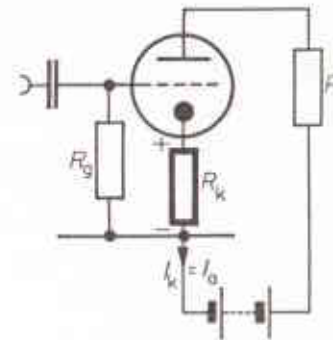


Abb. 6.11

An sogenannten Katodenwiderstand R_k der Schaltung tritt ein Spannungsabfall $U_k = I_a \cdot R_k$ auf, dessen positives Potential an der Katode liegt. Über einen hochohmigen Widerstand R_g (ca. $1 \text{ M}\Omega$) wird das negative Potential (Masse) an das Steuergitter gelegt. Somit ist das Gitter immer negativ gegenüber der Katode. Da sich Schwankungen des Anodenstroms unmittelbar auf die Gitterspannung auswirken und eine Gegenkopplung hervorrufen, kann der Katodenwiderstand durch einen Kondensator wechselstrommäßig überbrückt werden. Die Kapazität richtet sich nach der tiefsten zu übertragenden Frequenz.

Vorteile: Geringer Schaltungsaufwand. Schutzwirkung, da beim Ansteigen des Anodenstroms die Vorspannung des Steuergitters negativer wird.

Nachteile: Vorspannung belastungsabhängig. Katode liegt nicht auf Nullpotential, was in netzbetriebenen Verstärkern Brummen verursachen kann (vor allem in Vorverstärkern).

c) Einschalten eines sehr hochohmigen Widerstands zwischen Gitter und Katode

Hierbei nutzt man die Tatsache aus, daß in einer Elektronenröhre bei eingeschaltetem Heizsystem infolge der Elektronenemission bereits ohne äußere Spannung ein kleiner Anlaufstrom fließt. Dieser Anlauf-

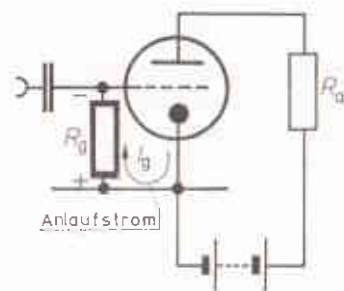


Abb. 6.12

$$R_g = \frac{U_g}{I_g} = \frac{0,75 \text{ V}}{50 \cdot 10^{-9} \text{ A}} = 15 \cdot 10^6 \Omega = 15 \text{ M}\Omega \text{ liegen.}$$

Die technische Stromrichtung des Anlaufstroms verläuft vom Gitter zur Katode, d.h. außerhalb der Röhre von der Katode zum Gitter. Somit liegt das negative Potential des Spannungsabfalls am Gitter.

Vorteile: Geringer Schaltungsaufwand. Keine Gegenkopplung. Verbindung Katode-Masse möglich. Daher geringe Brummneigung.

Nachteile: Anlaufstromstärke unterliegt starken Streuungen. Daher keine sehr exakte Arbeitspunkteinstellung. Im allgemeinen nur für Eingangsstufen mit geringer Steuerspannung anwendbar.

6.2.2. Verwendung der Röhre mit Steuergitter

Die Elektronenröhre mit Steuergitter findet in allen Gebieten der Elektrotechnik als Verstärkerröhre Verwendung:

- Verstärkung von Tonfrequenzen (Niederfrequenzverstärker), z.B. Fernsprechen, Tonbandgerät, Phonoverstärker.
- Verstärkung von Hochfrequenzen (Hochfrequenzverstärker), z.B. Trägerfrequenz-Fernsprechen, Rundfunk usw.

Daneben findet man sie z.B. als wichtiges Bauelement in Schaltungen zur Steuerung und Regelung (Rechenanlagen, automatische Steuerung von Maschinen usw.). Außer der Dreipolröhre gibt es heute für die verschiedensten Anwendungsgebiete eine Vielzahl von Röhrentypen mit bis zu sieben Gittern.

7. RC-Glieder

Zur Datenübermittlung werden in der Elektronik nahezu ausschließlich elektrische Impulse benutzt. Die Impulse werden jedoch auf dem Übertragungsweg durch Bauelemente wie Kondensatoren, Spulen, Dioden, Transistoren und Leitungen verzerrt. Die Verzerrung kann so weit führen, daß eine fehlerfreie Datenübertragung nicht gewährleistet ist. Ebenso treten häufig schon bei der Erzeugung von Impulsen bestimmter Form Schwierigkeiten auf. Die an die Erzeugung und Übertragung bestimmter Impulsformen gestellten hohen Anforderungen zwingen zu Schaltungsmaßnahmen, die eine Impulskorrektur bewirken. Am häufigsten sind Kombinationen von Widerständen (R) und Kapazitäten (C) in elektronischen Schaltungen Ursache für die Veränderung von Impulsformen. Die Wirkung solcher sog. RC-Glieder auf die Verzerrung und Entzerrung von elektrischen Signalen soll an einigen Beispielen aufgezeigt werden.

Induktivitäten (Spulen) sind in modernen Baugruppen der Elektronik so gut wie gar nicht anzutreffen. Das hängt u.a. damit zusammen, daß die Herstellung von Spulen in sog. integrierten Schaltungen oder auf Platinen für gedruckte Schaltungen schwierig oder gar unmöglich ist. Induktivitäten werden daher durch Baugruppen aus Widerständen, Kapazitäten und Transistoren nachgebildet.

7.1. RC-Glieder als Impulsformer

7.1.1. Elektrische Vorgänge beim Auf- und Entladen einer Kapazität

Bevor wir uns mit den eigentlichen RC-Gliedern beschäftigen, soll noch kurz auf das Verhalten eines Kondensators beim Anlegen einer Spannung und dem anschließenden Entladen eingegangen werden. Die grundsätzliche Wirkungsweise einer Kapazität beruht darauf, daß beim Anlegen oder Ändern einer Spannung im Dielektrikum eine begrenzte Elektronenverschiebung stattfindet. Die Elektronen können in Isolatoren ihre Schalen nicht verlassen. Daher ist nur eine Elektronenverschiebung innerhalb der Atome des Dielektrikums möglich. Sie wird als dielektrische Verschiebung bezeichnet. Beim Anlegen einer Spannung an die Beläge eines Kondensators dauert die dielektrische Verschiebung nur so lange, bis sich ein Gleichgewicht einstellt zwischen der elektrostatischen Kraftwirkung des äußeren elektrischen Feldes und der Kraft, mit der die Elektronen an den Atomkern gebunden sind. Wird die angelegte Spannung verändert, so ändert sich mit ihr die elektrostatische Kraftwirkung und die Elektronen werden erneut verschoben.

Wir können also folgern: Ändert man die Spannung an den Belägen eines Kondensators, so findet in seinem Dielektrikum eine sog. dielektrische Verschiebung statt. Da sich bei der dielektrischen Verschiebung Ladungsträger gerichtet bewegen, kann dieser Vorgang als elektrischer Strom bezeichnet werden. Solange die Spannung an den Kondensatorbelägen aber konstant bleibt, findet keine dielektrische Verschiebung statt; es fließt also kein Strom.

Wird eine Spannung mit Hilfe eines Schalters an den Kondensator gelegt, so wirkt er im ersten Augenblick wie ein Kurzschluß. Der dabei fließende Aufladungsstrom wird nur durch den inneren Widerstand der Spannungsquelle begrenzt. Im weiteren zeitlichen Verlauf steigt die Spannung am Kondensator bis auf die Höhe der EMK der Spannungsquelle. Dann ist die dielektrische Verschiebung beendet, es fließt kein Strom mehr. Trennt man die äußere Spannungsquelle wieder vom Kondensator, so behält er die Ladung so lange, bis er durch Anschalten eines Verbrauchers entladen wird.

Der zeitliche Verlauf der Lade- und Entladespannung sowie des Lade- und Entladestroms folgt einer sogenannten e-Funktion (= Kurve des natürlichen Wachstums).

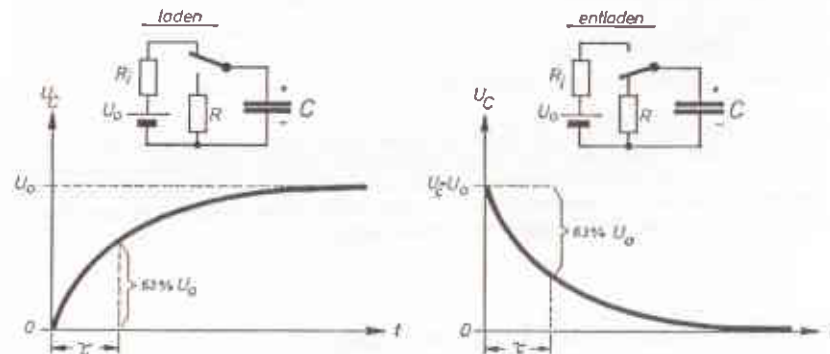


Abb. 7.1 — Zeitlicher Verlauf der Lade- und Entladespannung an einer Kapazität

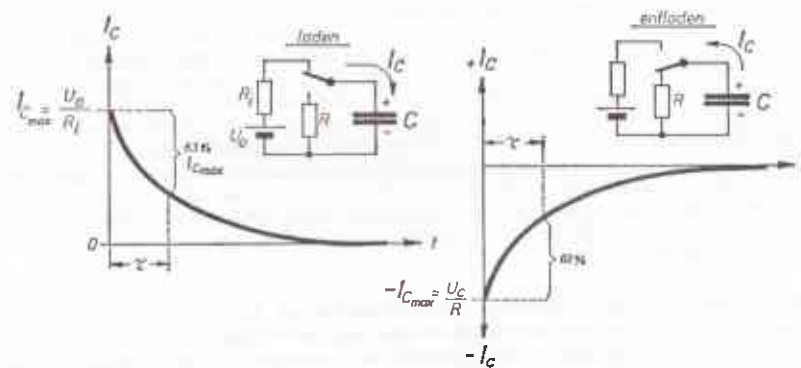


Abb. 7.2 — Zeitlicher Verlauf des Lade- und Entladestroms an einer Kapazität

Je größer die Kapazität (Speichervermögen) des Kondensators und der Widerstand der Spannungsquelle bzw. des Verbrauchers sind, desto länger dauert die Ladung bzw. Entladung. Als Richtwert für den zeitlichen Verlauf des Lade- und Entladevorgangs dient die **Zeitkonstante τ** ¹⁾. Sie läßt sich nach folgender Formel errechnen:

$$\text{Zeitkonstante} \quad \boxed{\tau = R \cdot C} \quad \text{s,}$$

wobei sich τ in Sekunden ergibt, wenn C in F und R in Ω eingesetzt werden.

Für den Umgang mit e-Funktionen müssen gute mathematische Kenntnisse vorausgesetzt werden. Ohne auf mathematische Einzelheiten einzugehen, kann jedoch folgende Angabe gemacht werden:

Spannung und Strom eines RC-Glieds haben sich nach der Zeit τ um 63 % und nach der Zeit 5τ um 99 % dem Endwert genähert (vgl. hierzu die Diagramme in Abb. 7.1 bis 7.3).

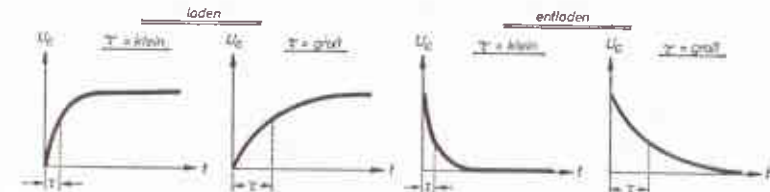


Abb. 7.3 — Zeitlicher Verlauf der Lade- und Entladespannung bei verschiedenen Zeitkonstanten

7.1.2. RC-Glied als Differenzglied

Der Einfachheit halber werden in den Beispielen zunächst wieder mechanische Schalter dargestellt. Das Ein- und Ausschalten wird in der Praxis durch elektrische Impulse vorgenommen.

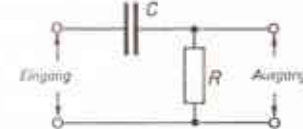


Abb. 7.4 — Differenzglied

Ein Differenzglied besteht aus einer **Kapazität im Längszweig** und einem **Wirkwiderstand im Querszweig** der Schaltung. Jede C-gekoppelte Transistor- oder Röhrenstufe stellt z.B. ein solches Differenzglied dar. Wir wollen jetzt untersuchen, wie sich die Ausgangsspannung an einem Differenzglied verhält, wenn eine Spannung an den Eingang geschaltet und anschließend der Eingang kurzgeschlossen wird.

Die Ausgangsspannung an einem Differenzglied verhält, wenn eine Spannung an den Eingang geschaltet und anschließend der Eingang kurzgeschlossen wird.

¹⁾ sprich Tau (griech. Buchstabe)

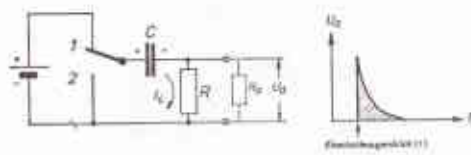


Abb. 7.5 — Einschalten der Eingangsspannung bei einem Differenzglied

gibt sich also ein Spannungsimpuls, dessen Form völlig unabhängig von der Einschaltdauer der angelegten Spannung ist. Sobald sich nämlich der Kondensator auf die Höhe der angelegten Spannung aufgeladen hat, fließt kein Strom mehr über R ; die Ausgangsspannung ist Null geworden.

Nach dem Umschalten des Schalters in Stellung 2 dient der geladene Kondensator als Spannungsquelle für den Stromkreis $C - R$. Beachtet man die Potentialverhältnisse, so zeigt sich, daß sich die Potentiale umgekehrt haben. Das Pluspotential der Kondensatorspannung liegt jetzt unten. Folglich fließt der Entladestrom in entgegengesetzter Richtung zum Ladestrom. An R tritt ein negativer Spannungsimpuls auf. Würden wir jetzt an den Eingang eines Differenzglieds eine Rechteckspannung legen, so würden sich die beschriebenen Vorgänge periodisch wiederholen. Das läßt sich sehr gut mit einem Oszilloskop darstellen. Dazu ist die Schaltung nach Abb. 7.7 aufzubauen.

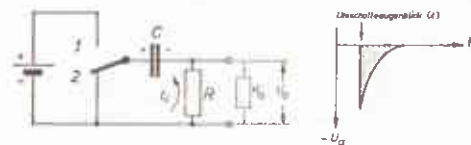


Abb. 7.6 — Kurzschließen des Eingangs eines Differenzglieds

negativer Spannungsimpuls auf. Würden wir jetzt an den Eingang eines Differenzglieds eine Rechteckspannung legen, so würden sich die beschriebenen Vorgänge periodisch wiederholen. Das läßt sich sehr gut mit einem Oszilloskop darstellen. Dazu ist die Schaltung nach Abb. 7.7 aufzubauen.

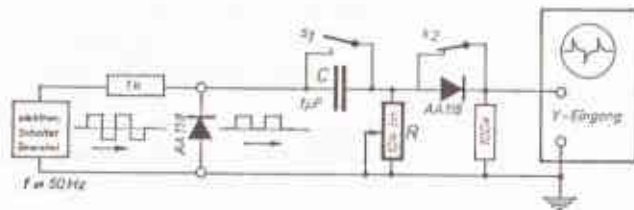


Abb. 7.7 — Differenzglied als Impulsformer

Als Generator für Rechteckspannungen dient ein elektronischer Schalter. Für die negativen Rechteckimpulse stellt die Germaniumdiode AA 118 nahezu einen Kurzschluß dar (Schalterstellung 2 in Abb. 7.6). Das Differenzglied besteht aus einem Kondensator mit $C = 1 \mu\text{F}$ und einem Regelwiderstand $R = 10\,000 \Omega$.

Als Zeitkonstante ergibt sich für diese Kombination

$$\tau = C \cdot R = 1 \cdot 10^{-6} \frac{\text{As}}{\text{V}} \cdot 10 \cdot 10^{-3} \frac{\text{V}}{\text{A}} = 10 \cdot 10^{-9} \text{ s} = 10 \text{ ns}.$$

Beträgt die Frequenz der Rechteckwechselspannung 50 Hz, so ergibt sich für die Impulsdauer

$$T_i \approx \frac{1}{2f} = \frac{1}{2 \cdot 50} \text{ s} = 1 \cdot 10^{-2} \text{ s}, \text{ also ebenfalls } 10 \text{ ms}.$$

Bei geschlossenem Schalter s_1 zeigt das Oszillogramm die zugeführten Rechteckimpulse. Wird der Schalter geöffnet, ist die Impulsformung durch das Differenzglied erkennbar. Auf dem Bildschirm zeigen sich positive und negative spitze Spannungsimpulse. Die Form der Impulse kann jetzt durch Verändern des Widerstands R beeinflusst werden. Je kleiner R wird, desto schmaler werden die Spannungsimpulse; denn mit R verringert sich die Zeitkonstante des RC-Glieds. Schaltet man hinter ein Differenzglied eine Diode, so läßt diese je nach Durchlaßrichtung nur die positiven oder negativen Impulse durch. Die Diode hat hier eine Schalterfunktion. Im Versuchsaufbau nach Abb. 7.7 kann diese Diode durch Öffnen des Schalters s_1 in den Stromkreis gelegt werden. Da der Eingangswiderstand eines Oszilloskops sehr groß ist, muß die Diode durch einen zusätzlichen Widerstand (ca. 100 k Ω) belastet werden.

Ein Differenzglied formt also Impulse von größerer Impulsdauer in kurzzeitige Impulse um. Die Impulsdauer ist um so geringer, je kleiner die Zeitkonstante des RC-Glieds ist.

7.1.3. RC-Glied als Integrierglied

Beim Integrierglied liegt der Widerstand eines RC-Glieds im Längszweig, die Kapazität im Querszweig der Schaltung. Mit einem Integrierglied läßt sich die Dauer der zugeführten Impulse vergrößern. Zur Erläuterung bedienen wir uns wieder der vereinfachten Schaltung, bei der die Steuerung durch einen mechanischen Schalter vorgenommen wird. Bei dieser Schaltung wird der Eingang des RC-Glieds für die Entladung des Kondensators unterbrochen. Da nämlich der Verbraucherwiderstand R parallel zum Kondensator liegt, kann dieser sich direkt über ihn entladen.

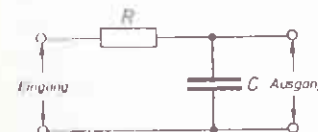


Abb. 7.8 — Integrierglied

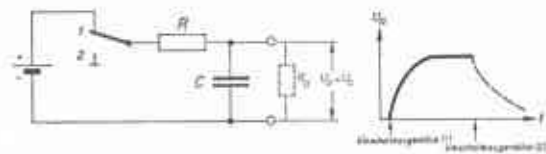


Abb. 7.9 — Einschalten der Eingangsspannung bei einem Integrierglied

In Schalterstellung 1 wird der Kondensator C über den Widerstand R aufgeladen, wobei der zeitliche Verlauf des Ladestroms wieder von der Zeitkonstante des RC-Glieds abhängt. Nach Beendigung des Ladevorgangs fließt kein Ladestrom mehr. Am Kondensator liegt eine bestimmte konstante Spannung. Wird der Schalter in Stellung 2 umgeschaltet, so kann sich der Kondensator über den Verbraucherwiderstand R_a entladen. Ist der Verbraucherwiderstand unendlich groß, bleibt die Ausgangsspannung konstant. Dagegen entlädt sich der Kondensator schnell, wenn R_a klein ist. Der zeitliche Verlauf des Entladestroms bzw. der Entladespannung wird hier also von der Größe des Verbraucherwiderstands wesentlich mitbestimmt ($\tau_E = C \cdot R_a$).

Legt man an den Eingang eines Integrierglieds Rechteckimpulse, so tritt abhängig von der Pulsfrequenz ein periodischer Wechsel zwischen Ladung und Entladung des Kondensators auf. Das läßt sich wieder am besten mit Hilfe eines Oszilloskops veranschaulichen (Schaltung nach Abb. 7.10).

Die Diode AA 118 ist für die negativen Rechteckimpulse als geöffneter Schalter anzusehen (vgl. Schalterstellung 2 in Abb. 7.9). Bei geöffnetem Schalter s zeigt das Oszillogramm die zugeführten Rechteck-

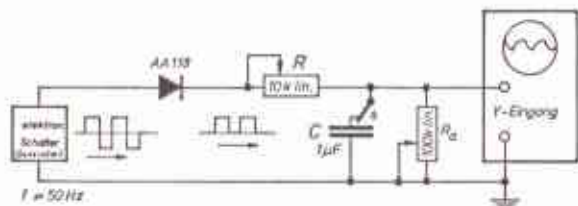
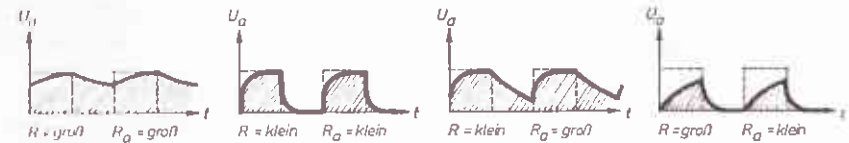


Abb. 7.10 — Impulsformung durch ein Integrierglied

impulse. Wird der Kondensator durch Schließen des Schalters in die Schaltung eingefügt, so wird die Impulsformung durch das Integrierglied deutlich. Durch Verändern von R läßt sich nun der Verlauf der Anstiegsflanken, durch Verändern von R_a der Verlauf der Abfallflanken beeinflussen.

Abb. 7.11 — Abhängigkeit der Ausgangsspannung von R und R_a

Neben ihrem Hauptanwendungsgebiet in der elektronischen Datenverarbeitung werden Integrierglieder z.B. in der **Stromversorgungstechnik** zur Glättung der aus Gleichrichtern gewonnenen pulsierenden Gleichspannung verwendet (Siebketten). Eine für die Meßtechnik interessante Anwendung ist der **Spitzenspannungsmesser**, wie er beim Rundfunk und in Tonbandgeräten zur Aussteuerungsüberwachung dient. Ein Spitzen-

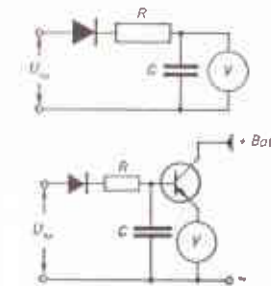


Abb. 7.12 — Spitzenspannungsmesser

Ist der Widerstand des Zeigerinstruments zu klein, so schaltet man einfach eine Transistorstufe in Kollektorschaltung davor (vgl. Kollektorschaltung ergibt hohen Eingangswiderstand und niedrigen Ausgangswiderstand).

Auch die **elektronischen Drehzahlmesser** für Kraftfahrzeugmotoren enthalten Integrierglieder, um eine Anzeige der Zündimpulse mit Zeigerinstrumenten zu ermöglichen. Bei Drehzahlmessern wird die Zahl der Zündimpulse in der Zeiteinheit gemessen. Der Mittelwert der Ausgangsspannung des Integrierglieds hängt dabei von der Häufigkeit der Zündimpulse in der Zeiteinheit ab.

Abschließend ist hierzu festzuhalten, daß ein Integrierglied die **Impulsdauer verlängert**. Die Zeitkonstante eines Integrierglieds kann im Verhältnis zur Impulsdauer der Impulsspannung so groß sein, daß mehrere kurzzeitig aufeinanderfolgende Impulse zu einem langen Impuls zusammengefaßt (= integriert) werden.

7.1.4. Einfache Kippschaltung mit Integrierglied

Schaltet man anstelle eines Widerstands R_a eine Glühlampe (ohne eingebauten Vorwiderstand!) hinter ein Integrierglied, so lassen sich

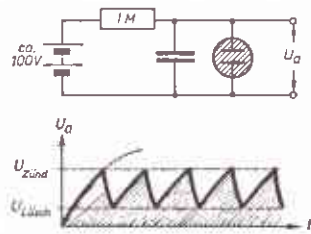


Abb. 7.13 — Erzeugung von Kippschwingungen

mit dieser Schaltung Kippschwingungen erzeugen. Am Eingang muß dazu eine Gleichspannung liegen, die mindestens die Höhe der Glühlampen-Zündspannung besitzt. Die Zeitkonstante für die Schaltung nach Abb. 7.13 ist mit $C = 1 \mu\text{F}$ so groß, daß die Kippfrequenz an der Glühlampe mit bloßem Auge erkennbar wird. Soll der Verlauf der Kippspannung im Oszillogramm dargestellt werden, so wird zweckmäßig eine Kapazität von 10 nF eingesetzt.

Nach dem Einschalten der Spannung lädt sich der Kondensator über R langsam auf, bis die Zündspannung für die Glühlampe (ca. 70 V) erreicht wird. Mit dem Zünden wird die Glühlampe niederohmig. Der Kondensator wird entladen, bis die Löschspannung der Glühlampe (ca. 40 V) unterschritten wird. In diesem Augenblick erlischt die Glühlampe. Ihr Widerstand wird wieder unendlich groß. Der Vorgang wiederholt sich periodisch. Die Kippfrequenz wird dabei wesentlich von der Höhe der angelegten Spannung und der Zeitkonstante des Integrierglieds bestimmt.

7.1.5. RC-Glied als Beschleuniger

Beim Schließen eines induktiv belasteten Stromkreises erreicht der Strom erst nach einer gewissen Anstiegszeit seine volle Stärke. Ebenso steigt die Spannung an einer Kapazität nach dem Anlegen einer Spannung erst nach einer bestimmten Zeit auf den Höchstwert. Dieses Verhalten von Induktivität und Kapazität führt bei der Übertragung von Impulsen mit steilen Anstiegsflanken — z.B. Rechteckimpulsen —

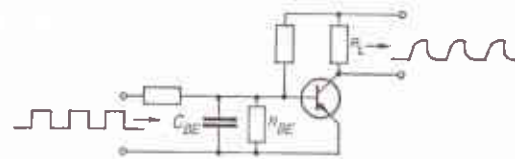


Abb. 7.14 — Verzerrung von Impulsen durch einen Transistor

zu Verzerrungen. Allein die Basis-Emitter-Kapazität eines Transistors wirkt sich schon nachteilig aus. Je höher die Impulsfrequenz und je größer die Kapazität ist, desto stärker wirken sich die Verzerrungen auf die Impulsform aus. Umschaltvorgänge in elektronischen Schaltungen (z.B. Multivibratoren oder Transistorschalter) müssen aber durch Spannungs- oder Stromstöße mit steiler Anstiegsflanke vorgenommen werden.

Zur Korrektur der im wesentlichen durch die Eingangskapazitäten von Transistoren verursachten Verzerrungen werden RC-Glieder als sogenannte Beschleuniger eingesetzt. Ein Beschleuniger besteht aus der Parallelschaltung eines Widerstands mit einem Kondensator und wird in Reihe zum Verbraucher geschaltet. Ein einfaches Beispiel ist in Abb. 7.15 wiedergegeben. Im Einschalt Augenblick hat der Kondensator den Widerstand 0 Ohm . Der Einschaltstrom wird nur durch R_a begrenzt. Nach dem Einschalten lädt sich der Kondensator auf — sein Widerstand wird mit zunehmender Aufladung unendlich groß. Der Gleichstrom wird nun durch $R_a + R$ begrenzt und ist kleiner als im Einschalt Augenblick. Die hohe Einschaltspitze eines Beschleunigers dient bei der Ansteuerung einer Transistorstufe dazu, die Verzögerung des Spannungsanstiegs zu kompensieren. Wird eine Transistorstufe über einen vorgeschalteten Beschleuniger mit Rechteckimpulsen angesteuert, so ist praktisch keine Verzerrung der Anstiegsflanke mehr nachweisbar. Betrachtet man die Ansteuerung des Transistors im zeitlichen Verlauf, so ergibt sich ein Übergang von Spannungssteuerung (R durch C kurzgeschlossen) zu Stromsteuerung (R voll wirksam als Vorwiderstand).

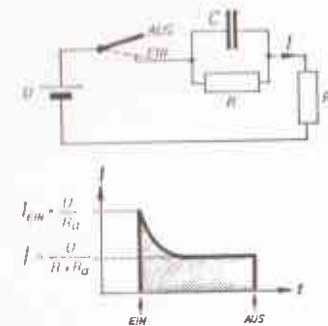


Abb. 7.15 — RC-Glied als Beschleuniger

Wird eine Transistorstufe über einen vorgeschalteten Beschleuniger mit Rechteckimpulsen angesteuert, so ist praktisch keine Verzerrung der Anstiegsflanke mehr nachweisbar. Betrachtet man die Ansteuerung des Transistors im zeitlichen Verlauf, so ergibt sich ein Übergang von Spannungssteuerung (R durch C kurzgeschlossen) zu Stromsteuerung (R voll wirksam als Vorwiderstand).

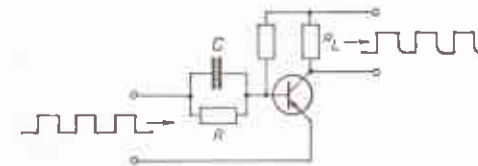


Abb. 7.16 — Transistorstufe mit Beschleuniger

7.2. RC-Glieder als Hoch- und Tiefpaß

RC-Glieder haben nicht nur als Impulsformer, sondern auch in der allgemeinen Übertragungstechnik große Bedeutung; denn sie bevorzugen oder benachteiligen bei der Übertragung eines breiten Frequenzbandes bestimmte Frequenzbereiche und rufen außerdem noch Phasenverschiebungen hervor. Da jede Schaltung zwangsläufig eine Vielzahl von Kapazitäten und Widerständen enthält, tritt auch — beabsichtigt oder unbeabsichtigt — bei der Übertragung eines Frequenzbandes eine Frequenzbeeinflussung auf. Je nach Kombination von Widerstand und Kapazität haben die RC-Glieder die Wirkungsweise eines Hoch- oder Tiefpasses.

7.2.1. Differenzierglied als Hochpaß

Ein Differenzierglied hat die Wirkungsweise eines Hochpasses. Legt man an seinen Eingang ein Frequenzgemisch — z.B. das Tonfrequenzspektrum von Sprache oder Musik —, so werden die hohen Frequenzen bevorzugt durchgelassen. Man kann das Differenzierglied als Wechselspannungsteiler ansehen. Da der kapazitive Blindwiderstand nach der Formel

$$X_C = \frac{1}{\omega \cdot C}$$

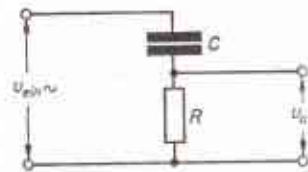


Abb. 7.17 — Differenzierglied als Wechselspannungsteiler

frequenzabhängig ist, ändert sich mit der Frequenz auch das Verhältnis X_C zu R . X_C wird mit zunehmender Frequenz kleiner, so daß sich die Teilspannungen zugunsten von U_a verschieben. Je höher also die Frequenz der zugeführten Wechselspannung ist, desto größer ist die Ausgangsspannung des Hochpasses.

Um eine eindeutige Aussage über die Wirkung einer Frequenzeinflussung machen zu können, hat man in der Übertragungstechnik den Begriff Grenzfrequenz eingeführt.

Als Grenzfrequenz f_0 eines RC-Glieds gilt die Frequenz, bei der $X_C = R$ ist. Unter dieser Bedingung liegt am Ausgang des RC-Glieds noch 70,7 % der Eingangswechselspannung.

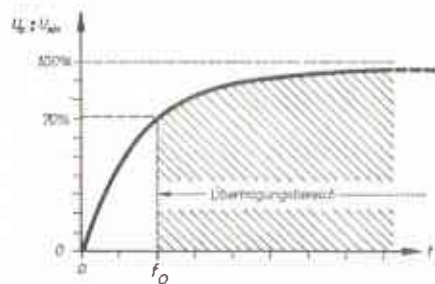


Abb. 7.18 — Ausgangsspannung eines RC-Hochpasses in Abhängigkeit von der Frequenz

Bei einem Hochpaß ist damit die untere Frequenzgrenze festgelegt. Die Wirkungsweise eines RC-Hochpasses wird durch nebenstehendes Diagramm verdeutlicht. Es wird meßtechnisch aufgenommen, indem man an den Eingang Wechselspannungen gleicher Höhe mit steigender Frequenz legt und zu jeder Frequenz die Ausgangsspannung mißt. Dabei zeigt sich, daß die Ausgangsspannung bis zur Grenzfrequenz zunächst stark, oberhalb der Grenzfrequenz jedoch nur noch wenig ansteigt.

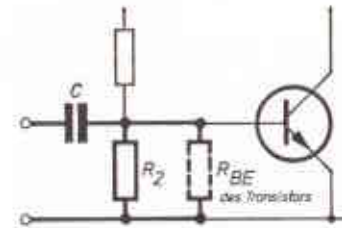


Abb. 7.19 — C-Kopplung als Differenzierglied

In Transistorverstärkern mit C-Kopplung wird durch den Koppelkondensator C und den Eingangswiderstand der Verstärkerstufe ($R_2 \parallel R_{BE}$) ein Differenzierglied gebildet. Wie schon bei der C-Kopplung angegeben, muß darauf geachtet werden, daß die Kapazität des Koppelkondensators groß genug ist. Nur damit ist sicherzustellen, daß tiefe Frequenzen ausreichend übertragen bzw. verstärkt werden.

Beispiel:

Zur C-Kopplung einer Transistorstufe liegt im Eingang ein Koppelkondensator mit $C = 1 \mu\text{F}$. Der Eingangswiderstand ($R_2 \parallel R_{BE}$) der Transistorstufe beträgt $1 \text{ k}\Omega$.

Wie groß ist die Grenzfrequenz?

Gegeben: $R = 1 \text{ k}\Omega$; $C = 1 \mu\text{F}$

Gesucht: f_0

Lösung: $X_C = R$

$$\frac{1}{\omega \cdot C} = R$$

$$\frac{1}{2 \pi \cdot f_0 \cdot C} = R$$

$$f_0 = \frac{1}{2 \pi \cdot R \cdot C} = \frac{1}{2 \pi \cdot 1 \cdot 10^3 \cdot 1 \cdot 10^{-6}} \text{ Hz}$$

$$f_0 = \frac{10^3}{2 \pi \cdot 1 \cdot 1} = \frac{10^3}{2 \pi} = \underline{\underline{159 \text{ Hz}}}$$

Die C-Kopplung überträgt brauchbar nur Frequenzen oberhalb 159 Hz.

7.2.2. Integrierglied als Tiefpaß

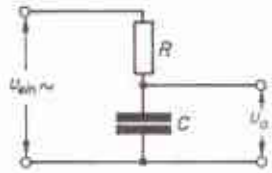


Abb. 7.20 — Integrierglied als
Wechselspan-
nungsteiler

Ein Integrierglied stellt einen Tiefpaß einfachster Form dar. Bei der Übertragung eines Frequenzgemisches werden von einem Integrierglied die tiefen Frequenzen bevorzugt zum Ausgang durchgelassen. Auch das Integrierglied stellt einen Wechselspannungsteiler dar. Hier liegt die Kapazität jedoch parallel zum Ausgang. Je höher die Frequenz der zugeführten Wechselspannung ist, desto kleiner ist der Blindwiderstand X_C des Kondensators. Er stellt also einen mit zunehmender Frequenz stärkeren Nebenschluß zum Ausgang dar. Der größere Spannungsabfall tritt dann am Widerstand R des Integrierglieds auf. Daraus folgt: Je größer die Frequenz der zugeführten Wechselspannung ist, desto kleiner ist die Ausgangsspannung.

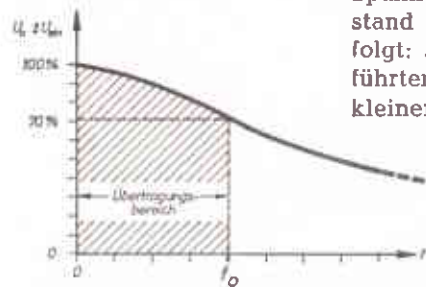


Abb. 7.21 — Ausgangsspannung eines
RC-Tiefpasses in Abhängig-
keit von der Frequenz

Als Grenzfrequenz f_0 gilt auch beim Integrierglied die Frequenz, bei der $X_C = R$ ist. Beim Tiefpaß wird damit die obere Frequenzgrenze festgelegt. Die Ausgangsspannung beträgt dabei noch 70,7 % der Eingangsspannung.

In Transistorverstärkern wird durch die Eingangskapazität C_{BE} des Transistors in Verbindung mit anderen, in Reihe liegenden Widerständen häufig ein Integrierglied gebildet. Je größer die Basis-Emitter-Kapazität des Transistors und der Reihenwiderstand sind, desto niedriger liegt die obere Frequenzgrenze der Verstärkerstufe. Soll also bei einer Transistorverstärkerstufe die obere Grenzfrequenz möglichst hoch liegen, so muß der vor dem Eingang liegende Widerstand so klein wie möglich gehalten werden.

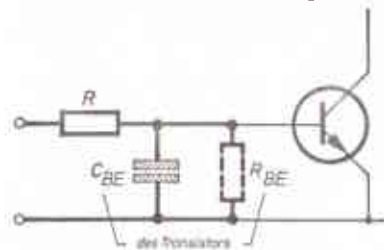


Abb. 7.22 — Eingang einer
Transistorstufe als
Integrierglied

Integrierglieder werden in der Übertragungstechnik — z.B. in NF-Verstärkern — dazu benutzt, den Frequenzgang gezielt zu beeinflussen. Ersetzt man den Widerstand R eines Integrierglieds durch einen Regelwiderstand, so können hohe Frequenzen ganz nach Belieben mehr oder weniger abgeschwächt werden. Solche regelbaren Integrierglieder werden als Tonblende bezeichnet.

Einfache aus Widerstand und Kapazität bestehende Siebglieder hinter Netzgleichrichtern haben ebenfalls Tiefpaßwirkung: Für den Welligkeitsanteil der pulsierenden Gleichspannung stellt der Blindwiderstand des Siebkondensators einen niederohmigen Nebenschluß dar, während der Gleichspannungsanteil ungehindert passieren kann.

8. Kippstufen

8.1. Allgemeines

Kippstufen sind wesentliche Grundbausteine in elektronischen Geräten. Man findet sie sowohl in einfachen Schaltungen der Digitaltechnik¹⁾ als auch in elektronischen Rechenanlagen. Sie spielen also in der Elektronik — insbesondere als Speicherelemente — eine erhebliche Rolle.

Kippstufen sind elektronische Schaltungen, die nur zwei bestimmte Schaltzustände einnehmen können. Sie bestehen grundsätzlich aus zwei Transistoren, die so extrem rückgekoppelt sind, daß beide nur als Schalter arbeiten können. Das bedeutet, die Transistoren werden entweder völlig durchgesteuert oder völlig gesperrt. Dabei wirkt die Rückkopplung so, daß der zweite Transistor zwangsläufig gesperrt ist, wenn der erste durchgeschaltet wird. Umgekehrt wird aber auch der andere durch den gesperrten Transistor zwangsweise leitend.

Je nach Art der Rückkopplung, die zwischen den beiden Transistoren besteht, ist zwischen drei verschiedenen Kippstufen zu unterscheiden:

a) Bistabile Kippstufe, auch Flipflop genannt:

Beide Transistoren sind hier galvanisch — über Widerstände — miteinander gekoppelt. Die Wirkungsweise einer bistabilen Kippstufe entspricht der eines polarisierten Relais (Telegraphenrelais).

b) Monostabile Kippstufe:

Die Kopplung zwischen den beiden Transistorstufen ist einmal galvanisch und zum anderen kapazitiv. Eine monostabile Kippstufe arbeitet wie ein neutrales Relais (z.B. Flachrelais 48).

c) Astabile Kippstufe:

Beide Transistoren sind kapazitiv — über Kondensatoren — miteinander gekoppelt. Hier tritt während des Betriebes kein stabiler Schaltzustand auf. Die astabile Kippstufe wird als Impulsgenerator eingesetzt.

Auf die Eigenarten der drei verschiedenen Kippstufen soll anschließend näher eingegangen werden. Häufig ist anstelle der Bezeichnung Kippstufe der Ausdruck Multivibrator anzutreffen.

Als entscheidendes Wesensmerkmal für die Kippstufen können wir uns aber bereits jetzt schon merken: **Bei einer Kippstufe ist stets ein Transistor gesperrt und der andere ist leitend.** Die verschiedenen Kippstufen können sowohl mit PNP- als auch mit NPN-Transistoren betrieben werden. In den folgenden Darstellungen wollen wir uns jedoch auf Schaltungen mit NPN-Transistoren beschränken. Als Betriebsspannung wählen wir $U = 9 \text{ V}$.

¹⁾ digital = sprunghaft, schrittweise

8.2. Signalzustände beim Transistor als Schalter

Da wir bereits den NPN-Transistor als Schalter kennengelernt haben, sind uns folgende Tatsachen bekannt: Führt man der Basis

— also dem Eingang — ein gegenüber dem Emittor positives Potential zu, so wird der NPN-Transistor leitend, d.h., die Kollektor-Emittorstrecke wird niederohmig. Sieht man von dem geringen Spannungsabfall ($U_{CEsat} \approx 0,2 \text{ V}$) an der Kollektor-Emittorstrecke ab, so führt der Kollektor (Ausgang) im durchgeschalteten Zustand Nullpotential (genau $+0,2 \text{ V}$, vgl. hierzu Abb. 8.1). Wird der Basis dagegen Nullpotential zugeführt, so sperrt der Transistor, d.h., die Kollektor-Emittorstrecke wird hochohmig. Sehen wir jetzt von dem geringfügigen Spannungsabfall an R_C — der durch den Sperrstrom verursacht wird — ab, so liegt die gesamte Spannung am Transistor. Der Ausgang A führt also jetzt das Potential von nahezu $+9 \text{ V}$ (vgl. hierzu Abb. 8.2). Daraus erkennen wir, daß der Transistor im Schalterbetrieb am Ausgang stets das entgegengesetzte Potential gegenüber dem Eingang

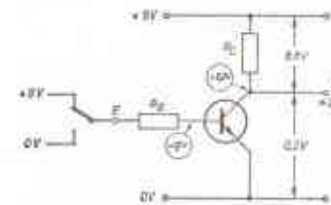


Abb. 8.1 — Potentiale bei durchgeschaltetem Transistor

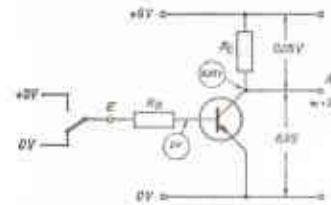


Abb. 8.2 — Potentiale beim gesperrten Transistor

führt. In der Elektronik wird in einem solchen Fall von einer **Signalumkehrung** gesprochen.

Den beiden möglichen Schaltzuständen am Ein- und Ausgang des Transistors können bestimmte durch Vereinbarung festgelegte Bedeutungen zugeordnet werden, z.B.

$+9 \text{ V}$ bedeutet Signal bzw.
 0 V bedeutet kein Signal.

Ebensogut kann man aber auch vereinbaren:

0 V bedeutet Signal bzw.
 $+9 \text{ V}$ bedeutet kein Signal.

In der Digitaltechnik werden die beiden möglichen Zustände wie folgt bezeichnet:

kein Signal $\longrightarrow$ 1
Signal $\longrightarrow$ 0

Je nach Vereinbarung kann also sein:

$$\left. \begin{array}{l} +9 \text{ V} = 1 \\ 0 \text{ V} = 0 \end{array} \right\} \text{ positive Logik}$$

oder:

$$\left. \begin{array}{l} 0 \text{ V} = 1 \\ +9 \text{ V} = 0 \end{array} \right\} \text{ negative Logik}$$

Bei Verwendung von PNP-Transistoren würde anstelle des Potentials von $+9 \text{ V}$ das Potential von -9 V treten. Selbstverständlich müssen für die beiden Signalzustände 0 und 1 die Potentiale nicht unbedingt 0 V bzw. $+9 \text{ V}$ (-9 V) betragen; sie können je nach Schaltung und Betriebsspannung auch andere Werte annehmen und auch in bestimmten Toleranzen schwanken.

Wir werden später noch einmal auf die Signalvereinbarungen zurückkommen und uns auf die negative Logik festlegen.

Wenn bei der Behandlung der folgenden Schaltungen von den Potentialen 0 V und $+9 \text{ V}$ die Rede ist, so sind die geringfügigen Spannungsabfälle an R_C im gesperrten Zustand bzw. im Transistor im leitenden Zustand jeweils vernachlässigt. Die exakten Potentialzustände werden wir nur dann berücksichtigen, wenn die Wirkungsweise einer bestimmten Schaltung genau untersucht werden soll.

8.3. Die bistabile Kippstufe (Flipflop)

8.3.1. Grundschialtung

Eine bistabile Kippstufe wird auch als Flipflop (FF) oder bistabiler Multivibrator (BMV) bezeichnet.

Zunächst einmal hängen wir einer im Schalterbetrieb arbeitenden Transistorstufe eine weitere Stufe an (Abb. 8.3).

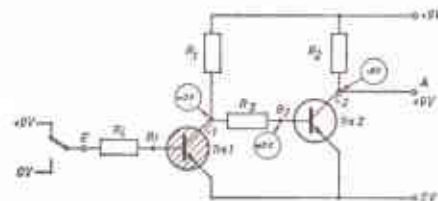


Abb. 8.3 — Zweistufiger Transistorschalter

Der Schaltung können wir sofort folgendes entnehmen:

1. Wenn Trs 1 leitend ist ($+9 \text{ V}$ am Eingang), dann beträgt das Potential an C_1 und damit an B_2 0 V . Der Transistor 2 ist also gesperrt. Der Ausgang (C_2) führt damit den Potentialzustand $+9 \text{ V}$.

2. Ist Trs 1 gesperrt (0 V am Eingang), dann ist das Potential an C_1 $+9 \text{ V}$. B_2 erhält damit positives Potential. Der Trs 2 wird leitend und schaltet 0 V zum Ausgang durch.

Wir können daraus folgende Schlußfolgerungen ziehen:

- a) Durch die Kopplung des Ausgangs (C_1) von Trs 1 mit dem Eingang (B_2) des Trs 2 ist der Schaltzustand des Trs 2 stets abhängig vom Schaltzustand des Trs 1.
- b) Die Kopplung ist derart, daß der Trs 2 stets den umgekehrten Schaltzustand von Trs 1 einnimmt.
- c) Das Potential am Eingang E und Ausgang A ist gleich.

Nun wollen wir folgenden Versuch machen: Wir stellen den Schalter am Eingang auf $+9 \text{ V}$; damit haben wir am Ausgang ebenfalls wieder $+9 \text{ V}$. Nun stellen wir zwischen dem Ausgang und dem Eingang eine Verbindung her. Das bedeutet, wir koppeln den Ausgang (C_2) des Trs 2 mit dem Eingang (B_1) des Trs 1 (vgl. hierzu Abb. 8.4).

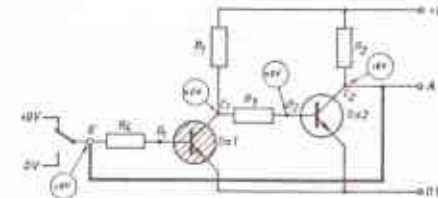


Abb. 8.4 — Zweistufiger Transistorschalter mit Rückkopplung vom Ausgang zum Eingang

Schalten wir jetzt die Spannung von $+9 \text{ V}$ am Eingang ab, so ändert sich am Zustand der beiden Transistoren nichts, denn am Eingang liegt nach wie vor über die Verbindung vom Ausgang zum Eingang das Potential von $+9 \text{ V}$. Das bedeutet, der Trs 1 wird durch das Ausgangspotential des Trs 2 in leitendem Zustand gehalten. Andererseits wird der Trs 2 durch den leitenden Zustand des Trs 1 weiterhin in gesperrtem Zustand gehalten. Die Schaltung ist also in dieser Lage stabil.

Machen wir nun noch einen 2. Versuch. Wir trennen die Verbindung Ausgang—Eingang wieder auf und stellen den Schalter am Eingang jetzt auf 0 V . Damit führt der Ausgang ebenfalls 0 V . Wir stellen die Verbindung Ausgang—Eingang wieder her und führen damit Nullpotential vom Ausgang zum Eingang. Nehmen wir jetzt das Nullpotential am Eingang des Schalters weg, so ändert sich am Schaltzustand der beiden Transistoren ebenfalls nichts und Trs 1 wird durch das Nullpotential vom Ausgang des Trs 2 weiterhin in gesperrtem Zustand gehalten. Trs 2 erhält — da Trs 1 gesperrt bleibt — nach wie vor positives Eingangspotential und bleibt somit leitend. Die Schaltung ist also auch in dieser Lage stabil; es lassen sich demnach hiermit zwei stabile Lagen herstellen.

Diese Schaltung stellt auch bereits die Grundschialtung der bistabilen Kippstufe dar. Wenn wir sie etwas umzeichnen, so kommen wir zu der allgemein üblichen Darstellungsweise (Abb. 8.5).

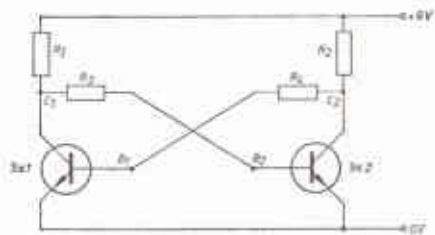


Abb. 8.5 — Grundsaltung der bistabilen Kippstufe

Bei dieser Darstellung wird besonders deutlich, daß jeweils der Ausgang des einen Transistors mit dem Eingang des anderen Transistors galvanisch gekoppelt ist.

R_1 und R_2 sind die Arbeitswiderstände (Lastwiderstände) für Trs 1 bzw. Trs 2 und R_3 und R_4 die Basisvorwiderstände für die beiden Transistoren.

Bistabile Kippstufen sind stets symmetrisch aufgebaut, d.h. $R_1 = R_2$ und $R_3 = R_4$. Bei völliger Symmetrie sämtlicher Bauelemente würden beim Einschalten der Betriebsspannung beide Transistoren nur halb durchgesteuert werden. In der Praxis ist jedoch eine völlige Übereinstimmung sämtlicher Bauelemente nie gegeben. Damit wird auch beim Einschalten der Betriebsspannung einer der beiden Transistoren früher leitend werden und somit den anderen Transistor sperren.

Bei der Untersuchung einer Schaltung können wir also stets davon ausgehen, daß ein Transistor zufällig leitend und dadurch der andere gesperrt ist. Wir wollen nun einmal annehmen, der Trs 1 sei leitend, der Trs 2 also gesperrt. Unter dieser Voraussetzung untersuchen wir die genauen Potentialverhältnisse an den verschiedenen Punkten der Schaltung. Die Ergebnisse sind aus den Angaben in der Abb. 8.6 ersichtlich.

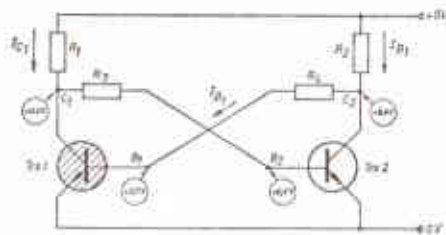


Abb. 8.6 — Potentialverhältnisse an einer bistabilen Kippstufe

Hinweis: In einer Schaltung stellen wir stets den als leitend angenommenen Transistor schraffiert dar.

Wie wir bereits wissen, liegt bei einem Silizium-Transistor die Diffusionsspannung der Basis-Emitter-Diode bei etwa 0,6 V. Um den Transistor völlig durchzusteuern, ist etwa eine Basisspannung von 0,7 V erforderlich. Dagegen beträgt die Kollektor-Emitter-Spannung im durchgeschalteten Zustand $U_{CEsat} \approx 0,2$ V. Das Potential an B_1 beträgt also +0,7 V, an $C_1 = +0,2$ V.

Damit liegt auch die Spannung an B_2 fest, denn sie kann nicht größer als +0,2 V sein. Da die Basis-Emitter-Strecke erst bei etwa +0,6 V leitend wird, muß der Trs 2 gesperrt bleiben.

Für die praktische Schaltungsausführung wollen wir den Silizium-NPN-Transistor BC 107 verwenden.

$$\begin{aligned} \text{Grenzdaten: } I_{Cmax} &= 100 \text{ mA} \\ P_{Cmax} &= 300 \text{ mW} \end{aligned}$$

Bei einer Eingangsspannung von $U_{BE} = 0,7$ V ist außerdem

$$\begin{aligned} I_B &= 50 \mu\text{A} \\ \beta &= 200 \quad (\text{Gleichstromverstärkung}). \end{aligned}$$

Diese Daten genügen, um die einzelnen Widerstände der Schaltung berechnen zu können. Der Kollektorstrom beträgt demnach $I_C = \beta \cdot I_B = 200 \cdot 50 \mu\text{A} = 10\,000 \mu\text{A} = 10 \text{ mA}$. Bei der Berechnung gehen wir davon aus, daß Trs 1 leitend ist. Da im durchgeschalteten Zustand am Transistor noch eine Restspannung von $U_{CEsat} \approx 0,2$ V liegt, muß an R_1 eine Spannung von $U_{R1} = 8,8$ V abfallen. Daraus läßt sich R_1 berechnen:

$$R_1 = \frac{U_{R1}}{I_{C1}} = \frac{8,8 \text{ V}}{0,01 \text{ A}} = \underline{\underline{880 \Omega}}$$

Wir wollen einen Widerstand von 820Ω — entsprechend der Normreihe für Widerstände — wählen. Da die Schaltung symmetrisch aufgebaut ist, gilt:

$$R_1 = R_2 = \underline{\underline{820 \Omega}}$$

Damit wird I_C etwas größer als 10 mA.

Das ist bei einem Transistor als Schalter sogar von Vorteil. Die Stromerhöhung führt nicht zur Überlastung des Transistors, solange

$$\begin{aligned} I_C &< 100 \text{ mA und} \\ P_C &< 300 \text{ mW ist.} \end{aligned}$$

Der tatsächliche Kollektorstrom beträgt:

$$I_C = \frac{U_{R1}}{R_1} = \frac{8,8 \text{ V}}{820 \Omega} = 0,0107 \text{ A} = \underline{\underline{10,7 \text{ mA}}}.$$

Zur Kontrolle wollen wir die Verlustleistung noch berechnen. Es ist:

$$P_C = I_C \cdot U_{CEsat} = 10,7 \text{ mA} \cdot 0,2 \text{ V} = 0,0107 \text{ A} \cdot 0,2 \text{ V} = \underline{\underline{0,00214 \text{ W}}}.$$

Damit liegen wir noch weit unter der maximal zulässigen Verlustleistung von 0,3 W.

Nachdem I_C feststeht, läßt sich der Mindestbasisstrom berechnen, der erforderlich ist, damit der Transistor den kleinstmöglichen Widerstand hat. Es ist:

$$I_{B1} = \frac{I_C}{\beta} = \frac{10,7 \text{ mA}}{200} = 0,0535 \text{ mA} = \underline{\underline{53,5 \mu\text{A}}}.$$

Da R_2 bekannt ist und über R_2 der Basisstrom I_{B1} fließen muß, können wir den Spannungsabfall R_2 errechnen und somit den Potentialzustand an C_2 bestimmen. Es ist:

$$U_{R2} = I_{B1} \cdot R_2 = 53,5 \mu\text{A} \cdot 820 \Omega = 53,5 \cdot 10^{-6} \text{ A} \cdot 820 \frac{\text{V}}{\text{A}}$$

$$U_{R2} = 0,044 \text{ V} \approx 0,05 \text{ V}$$

Das Potential an C_2 beträgt damit $+8,95 \text{ V}$.

Nun wollen wir R_4 bestimmen. Die Potentialdifferenz zwischen C_2 und B_1 muß an R_4 abfallen. Es ist:

$$R_4 = \frac{U_{R4}}{I_{B1}} = \frac{U_{C2} - U_{BE}}{I_{B1}} = \frac{8,95 \text{ V} - 0,7 \text{ V}}{0,0535 \cdot 10^{-3} \text{ A}}$$

$$= \frac{8,25}{0,0535} \cdot 10^3 \Omega = 154\,000 \Omega = 154 \text{ k}\Omega.$$

Wir wählen entsprechend der Normreihe mit Rücksicht auf eine kräftige Durchsteuerung der Transistoren den nächstkleineren Normwert:

$$R_3 = R_4 = 130 \text{ k}\Omega.$$

Damit ist der Basisstrom höher, nämlich:

$$I_{B1} = \frac{U_{R4}}{R_4} = \frac{8,25 \text{ V}}{130\,000 \Omega} = 0,0000635 \text{ A} = 63,5 \mu\text{A}.$$

Dies ist zulässig. Um den Trs voll durchzusteuern, darf I_B zwar größer, nicht aber kleiner als der berechnete Wert sein. Demnach müssen R_3 und R_4 bei der Rundung kleiner, nicht aber größer als berechnet, gewählt werden. Wir haben damit die Größen sämtlicher Widerstände bestimmt und kennen die Potentialverhältnisse an den Ein- und Ausgängen der Transistoren.

Zur Berechnung der Basisvorwiderstände bedient man sich für Schalttransistoren meistens der Erfahrungsformel

$$R_B = 0,8 \cdot R_L \cdot \beta$$

In unserem Fall ergibt sich

$$R_B = 0,8 \cdot 820 \Omega \cdot 200 = 131\,000 \Omega = 131 \text{ k}\Omega,$$

also fast genau der von uns gewählte Widerstand.

Bei künftigen Betrachtungen soll der Einfachheit halber das Potential am Kollektor des gesperrten Transistors ($+8,95 \text{ V}$) mit $+8,9 \text{ V}$ angenommen werden. Wäre Trs 2 leitend, so würden wir dieselben Potentialverhältnisse nur seitenvertauscht vorfinden.

8.3.2. Grundschialtung mit Hilfsspannung

Bei einer Eingangsspannung von $+0,2 \text{ V}$ ist ein Silizium-Transistor zwar noch sicher gesperrt, der Kollektorreststrom (Sperrstrom) hat aber damit noch nicht sein Minimum erreicht. Bei höheren Betriebstemperaturen (Sperrstrom steigt) oder bei Verwendung von Germanium-Transistoren — hier liegt der Sperrstrom höher und die Diffusionsspannung nur bei etwa $0,2 \text{ V}$ — empfiehlt es sich, an die Basis eine zusätzliche Sperrspannung anzulegen. In Abb. 8.7 wurde eine Hilfsspannung von -9 V gewählt. Die Widerstände R_5 und R_6 wählt man so, daß an der Basis des gesperrten Transistors etwa ein Potential von $-0,5 \text{ V}$ liegt.

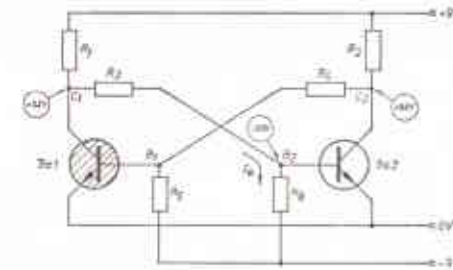


Abb. 8.7 — Bistabile Kippstufe mit Hilfsspannung

Selbstverständlich kann als Hilfsspannung auch eine andere Spannung als -9 V gewählt werden; dann müssen aber auch die Widerstände R_5 und R_6 entsprechend bemessen werden.

Wir wollen nun $R_5 = R_6$ berechnen. Zu diesem Zweck muß der Strom I_Q bestimmt werden, der über R_3 und R_6 fließt. Da R_3 und der Spannungsabfall an R_3 (Potentialdifferenz zwischen den Punkten C_1 und B_2) bekannt sind, können wir I_Q berechnen. Es ist:

$$I_Q = \frac{U_{C1} - U_{B2}}{R_3} = \frac{+0,2 \text{ V} - (-0,5 \text{ V})}{130\,000 \Omega}$$

$$= \frac{0,7 \text{ V}}{130\,000} = 0,0000054 \text{ A} = 5,4 \mu\text{A}.$$

Damit läßt sich dann auch R_6 berechnen. Es ist:

$$R_6 = \frac{U_{R6}}{I_Q} = \frac{8,5 \text{ V}}{5,4 \mu\text{A}} = \frac{8,5 \text{ V}}{5,4 \cdot 10^{-6} \text{ A}} = 1,57 \cdot 10^6 \Omega = 1,57 \text{ M}\Omega.$$

Wir wählen für

$$R_5 = R_6 = 1,5 \text{ M}\Omega.$$

Bei den künftigen Abbildungen wird jedoch aus Gründen der Übersichtlichkeit auf die Darstellung der Hilfsspannung verzichtet.

8.3.3. Bistabile Kippstufe mit statischen Eingängen

Die bistabile Kippstufe (Flipflop) gewinnt erst an Bedeutung, wenn wir ihren Zustand auch von außen her beeinflussen können. Bei der bisher besprochenen Grundschialtung war dies noch nicht möglich, da besondere Eingänge zur Steuerung des Flipflops fehlten. Abb. 8.8 zeigt ein Flipflop mit je einem Eingang zur Basis der beiden Transistoren (E_1 und E_2). Selbstverständlich wollen wir dann auch den Zustand des Flipflops irgendwo abgreifen; das geschieht an den Ausgängen A_1 und A_2 . Wir können uns schon jetzt merken: Ein Kippvorgang — das bedeutet, den leitenden Transistor sperren und den

gesperrten Transistor leitend machen — ist nur möglich, wenn einer der beiden Transistoren über einen besonderen Eingang von außen angesteuert wird.

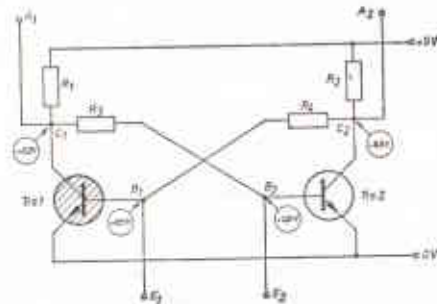


Abb. 8.8 — Bistabile Kippstufe mit Ein- und Ausgängen

Ein Kippvorgang kann sowohl über E_1 als auch über E_2 ausgelöst werden; er läßt sich erreichen durch:

- 0 V an der Basis des leitenden Transistors oder
- positives Potential (+0,7 V) an der Basis des gesperrten Transistors.

In Abb. 8.8. wurde Trs 1 als leitend angenommen. Ein Kippvorgang läßt sich also in diesem Fall erreichen durch:

- 0 V an E_1 oder
- +0,7 V an E_2 .

Im Fall a) wird der leitende Transistor durch 0 V an E_1 gesperrt; im Fall b) wird der gesperrte Transistor durch +0,7 V an E_2 leitend gemacht.

In digitalen Schaltkreisen treten im allgemeinen nur die beiden Potentialzustände Null oder volle Speisespannung auf. In unserem Falle also 0 V und +9 V (genauer +0,2 und +8,9 V). Der Eingang eines Flipflops wird daher von anderen elektronischen Bausteinen entweder mit 0 V oder mit +9 V angesteuert.

Ansteuerung mit positivem Potential

Wählt man die vorgenannte Möglichkeit b), nämlich das Ansteuern des gesperrten Transistors, so muß durch einen Vorwiderstand (R_E) dafür gesorgt werden, daß nicht die volle Spannung von +9 V an die Basis des Transistors gelangt. Durch die hohe positive Spannung wür-

de der Transistor zerstört werden. Der Vorwiderstand muß so gewählt werden, daß beim Anlegen einer Spannung von +9 V an den Eingang E an der Basis nur noch ein Potential von +0,7 V wirksam wird.

Um die Widerstände R_E berechnen zu können, gehen wir von den Potentialzuständen vor dem Kippvorgang aus (vgl. Abb. 8.8: +0,2 V an C_1 und B_2). Beim Anlegen der Steuerspannung von +9 V an E_2 kann — solange Trs 2 noch nicht durchgeschaltet ist — zunächst nur ein Strom I_Q von E_2 über B_2 nach C_1 fließen (siehe den in Abb. 8.9 hervorgehobenen Stromkreis).

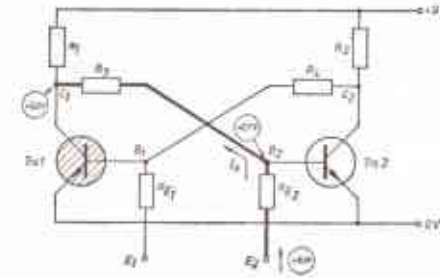


Abb. 8.9 — Potentialverhältnisse vor dem Kippvorgang beim Anlegen einer Durchschaltspannung

Durch diesen Strom I_Q muß an dem bereits bekannten Widerstand R_3 eine Spannung von

$$U_{R3} = U_{B2} - U_{C1} = 0,7 \text{ V} - 0,2 \text{ V} = 0,5 \text{ V}$$

abfallen, damit an B_2 ein Potential von + 0,7 V auftritt.

Da U_{R3} und R_3 bekannt sind, kann I_Q berechnet werden. Es ist

$$I_Q = \frac{U_{R3}}{R_3} = \frac{0,5 \text{ V}}{130 \text{ k}\Omega} = \frac{0,5 \text{ V}}{130\,000 \Omega} = 0,00000385 \text{ A} = 3,85 \mu\text{A}.$$

Da nun I_Q bekannt ist, kann auch R_E berechnet werden. Es ist:

$$\begin{aligned} R_E &= \frac{U_{RE}}{I_Q} = \frac{U_E - U_{B2}}{I_Q} = \frac{8,9 \text{ V} - 0,7 \text{ V}}{3,85 \cdot 10^{-6} \text{ A}} \\ &= \frac{8,2}{3,85} \cdot 10^6 \Omega = 2,13 \cdot 10^6 \Omega = 2,13 \text{ M}\Omega. \end{aligned}$$

Gewählt nach Normreihe: $R_E = 2 \text{ M}\Omega$.

Sobald Trs 2 durchgesteuert ist, wird $I_Q = 0 \text{ A}$, da Trs 1 jetzt gesperrt und somit an C_1 (+8,9 V) und an E_1 gleiches Potential liegt (Potentialdifferenz = 0 V, daher auch $I_Q = 0$). Das Potential an B_2 (+0,7 V) wird jetzt nur noch durch den Basisstrom bestimmt. Die entsprechende Berechnung haben wir ja bereits an früherer Stelle schon durchgeführt (siehe Berechnung von R_4).

Durch abwechselndes Anlegen von $+9\text{ V}$ an E_1 und E_2 kann man nun das Flipflop jeweils von der einen Lage in die andere Lage kippen. Da die Eingänge mit Gleichspannung angesteuert werden, spricht man bei diesen Eingängen von **statischen Eingängen**.

Ansteuerung mit Nullpotential

In der Praxis wird meistens ein Kippvorgang durch Ansteuerung mit Nullpotential ausgelöst. Das bedeutet, man steuert den leitenden Transistor mit Nullpotential an und sperrt ihn damit. Um dies zu erreichen, können wir Eingänge ohne Vorwiderstände wählen, da die Ansteuerung mit 0 V den Transistor nicht gefährden kann.

Wie aber schon an früherer Stelle erwähnt, sind Flipflops Bausteine elektronischer Anlagen. Das bedeutet, daß das Flipflop, je nach Zustand des vorgeordneten Bausteines, mit Nullpotential oder mit dem Potential der Betriebsspannung angesteuert werden kann. (In unserem Falle also mit 0 V oder $+9\text{ V}$.) Es muß aber jetzt in jedem Falle verhindert werden, daß auch $+9\text{ V}$ an die Basis des Transistors gelangen kann. Um dies zu erreichen, legen wir einfach zur Entkopplung Dioden in die Eingänge (vgl. hierzu Abb. 8.10).

Man muß dafür Ge-Dioden verwenden, da diese eine kleinere Diffusionsspannung ($0,2\text{ V}$) als Si-Dioden haben. Da an der Diode mindestens die Diffusionsspannung von $0,2\text{ V}$ abfällt, erhält man bei Ansteuerung des Eingangs E mit 0 V ein Potential von $+0,2\text{ V} \dots +0,3\text{ V}$ an der Basis des Transistors. Dieses Potential reicht aber noch aus, um einen Si-Transistor zu sperren. Bei Verwendung von Si-Dioden (Diffusionsspannung $0,6\text{ V}$) würde an der Basis ein Potential von $+0,6\text{ V} \dots +0,7\text{ V}$ liegen. Damit ist ein Sperren nicht mehr gewährleistet, da ein Si-Transistor mit einer Basis-Emitter-Spannung von etwa $0,6\text{ V}$ bereits durchgesteuert wird.

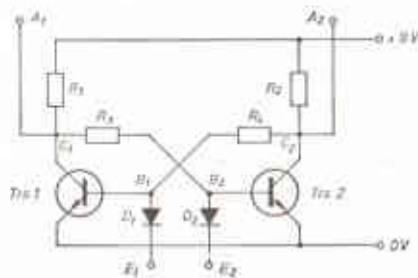


Abb. 8.10 — Statische Eingänge für Ansteuerung mit Nullpotential

Durch abwechselndes Anlegen von 0 V an die Eingänge E_1 und E_2 läßt sich auch in diesem Fall das Flipflop jeweils von der einen Lage in die andere Lage kippen und somit an den Ausgängen wechselweise 0 V oder $+9\text{ V}$ abgreifen.

Zusammenfassend ist demnach festzustellen, daß sich ein Flipflop durch Ansteuerung über die Eingänge in zwei stabile (daher „bistabil“) Schaltzustände überführen läßt. Je nach Schaltungsart der Eingänge kann ein Kippvorgang entweder durch Anlegen einer Steuerspannung oder von Nullpotential ausgelöst werden. Der Schaltzustand eines Flipflops wird durch die Potentialverhältnisse an den Ausgängen gekennzeichnet.

8.3.4. Bistabile Kippstufe mit dynamischen Eingängen

Im Abschn. 8.3.6 sind einige Anwendungsbeispiele für FF aufgeführt. Aus den Schaltsymbolen ist zu ersehen, daß die dort aufgeführten Schaltungen aus einzelnen FF bestehen, die über sogenannte dynamische Eingänge miteinander verbunden sind.

Wie wir aus Abschn. 8.3.3 wissen, nehmen FF zwei feste (stabile) Lagen ein und müssen von außen in die eine oder andere Lage umgeschaltet werden können. Dies kann durch eine positive Steuerspannung ($+0,7\text{ V}$) oder eine negative Steuerspannung ($\approx 0\text{ V}$) geschehen, die an die Basis des gerade gesperrten (Trs wird dann leitend) oder des gerade leitenden (Trs wird dann gesperrt) Transistors gelegt wird. In den bisher beschriebenen Schaltungen mit statischen Eingängen geschieht dies mit einer Gleichspannung als Steuerspannung, die auch dann noch anliegt, wenn das FF bereits umgeschaltet hat und in der neuen Lage stabil liegen bleibt. Zum Umschalten des FF würde jedoch ein entsprechend gepulster Steuerimpuls ausreichen, der das FF zum Kippen bringt und nach dessen Abklingen das FF sich in dem jetzt herrschenden Schaltzustand selbst hält. Die Polarität der Steuerimpulse wird in der Praxis so gewählt, daß der leitende Transistor gesperrt wird; d.h. bei NPN-Transistoren wird negatives und bei PNP-Transistoren positives Potential als Sperrimpuls verwendet.

Betrachten wir nun die Schaltung (Abb. 8.11), so stellen wir fest, daß sie sich von der vorhergehenden Schaltung (Abb. 8.10) durch die in die Steuerleitungen eingefügten Kondensatoren C_1 und C_2 und die Widerstände R_5 und R_6 unterscheidet. Weiterhin sind in der Regel die Eingänge E_1 und E_2 miteinander verbunden. Es ist zu erkennen, daß die Steuereingänge nichts anderes als jeweils ein Differenzierglied — bestehend aus C_1 und R_5 bzw. C_2 und R_6 — darstellen.

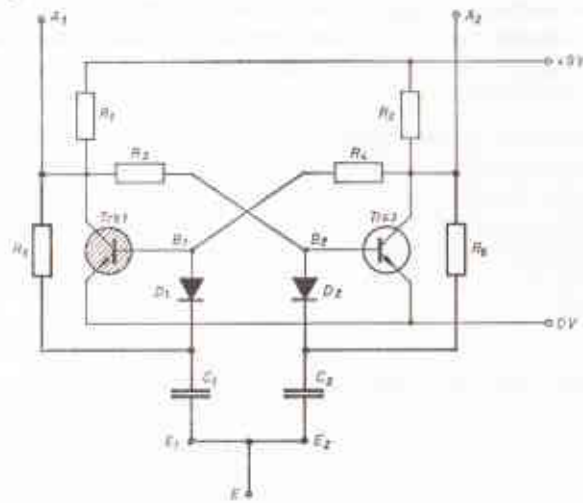


Abb. 8.11 — Bistabile Kippstufe mit dynamischen Eingängen

Die Differenzierglieder wandeln nun von außen angelegte Gleichspannungen in positive und negative Impulse um. Die Dioden D_1 und D_2 sperren die positiven Impulse, so daß nur negative Impulse an die Basen der beiden Transistoren gelangen können (vgl. hierzu auch Abschn. 7.1.2 („RC-Glied als Differenzierglied“)). Weiter fällt auf, daß die Widerstände R_5 und R_6 , die in einem üblichen Differenzierglied nur zur Entladung der Kondensatoren dienen, an die Ausgänge A_1 bzw. A_2 geschaltet sind. Mit Hilfe dieser Widerstände wird jeweils die an der Basis des gesperrten Transistors liegende Diode gesperrt, so daß auch kein negativer Impuls passieren kann, die andere Diode wird dagegen geöffnet.

Das geschieht auf folgende Weise: Aus Abb. 8.8 können wir ersehen, daß am Ausgang des leitenden Transistors $+0,2$ V und an seiner Basis $+0,7$ V, dagegen am Ausgang des gesperrten Transistors $+8,9$ V und an dessen Basis $+0,2$ V liegen. Die Diode D_1 hat (nach der Abb. 8.11) an der Anode $+0,7$ V, an der Katode $+0,2$ V; damit ist die Diode D_1 entsperrt. Die Diode D_2 hat dagegen an der Anode $+0,2$ V und an der Katode $+8,9$ V liegen, wodurch die Diode D_2 auch für negative Impulse gesperrt bleibt. Weil die Sperrimpulse nur den gerade leitenden Trs erreichen können, ist es ohne weiteres möglich, die Eingänge E_1 und E_2 parallelzuschalten.

8.3.5. Symbolische Darstellung des FF

Wir haben das Flipflop in allen Einzelheiten betrachtet. Es liegt auf der Hand, daß zur Darstellung umfangreicher Schaltungen, die eine Vielzahl solcher FF und anderer elektronischer Baugruppen enthalten, nicht mehr die Einzelstromlaufzeichnungen verwendet werden können; es müssen hier vielmehr symbolische Darstellungen gewählt werden.

Solche Schaltsymbole sind sicher aus der Wählvermittlungstechnik bekannt, wo für einen Hebdrehwähler oder für einen EMD-Wähler nur das entsprechende Symbol z.B. in einem Übersichtsplan verwendet wird. In der Abb. 8.12 ist das Symbol eines FF in ganz allgemeiner Darstellung gezeichnet. Hier ist ein rechteckiges Kästchen zu erkennen, das in der Mitte geteilt ist. Zur Unterseite führen zwei Eingänge (E_1 und E_2), nach oben zwei Ausgänge (A_1 und A_2). Vergleichen wir mit der Abb. 8.8, so ist zu erkennen, daß ein FF aus zwei Transistorschaltungen besteht, wobei jeweils einer der Transistoren leitet und dabei den anderen sperrt.

Wir wollen uns also merken, daß in jeder Kästchenhälfte ein Transistor verborgen ist, und daß jedem Transistor ein Ein- und ein Ausgang zugeordnet ist.

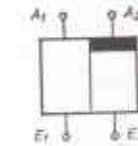


Abb. 8.13 —
FF mit fest-
gelegter
Ruhelage

Die Funktion vieler Schaltungen setzt voraus, daß ein FF immer eine bestimmte Ruhelage einnimmt (vgl. hierzu Abb. 8.13). Es handelt sich hierbei um das gleiche Symbol wie oben; es unterscheidet sich nur durch den schwarzen Balken auf der einen Seite. Dieser Balken bedeutet, daß der Ausgang auf der Seite des Balkens im Betriebszustand des FF ohne Ansteuerung das „Signal 1“ führt.

Nun müssen wir zunächst einmal überlegen, was mit den Begriffen Signal und 1 anzu-fangen ist. „Signal“ ist sicher nicht schwer zu verstehen. „Signal“ bedeutet Zeichen; es muß nur wieder eine Vereinbarung getroffen werden, was das Zeichen bedeuten soll. Da wir uns mit der Elektronik beschäftigen, müssen unsere Zeichen elektrischer Natur sein oder elektrisch erzeugt werden. So kann z.B. eine Glühlampe leuchten oder dunkel sein, ein Strom fließen oder nicht fließen, an einem bestimmten Punkt eine Spannung herrschen oder nicht. Bei diesen Zeichen handelt es sich also immer um sogenannte JA- oder NEIN-Entscheidungen; sie sind in der Regel sicher und unverwechselbar zu unterscheiden (denken Sie nur an die Glühlampe). Bewußt werden Unterscheidungen und Entscheidungen zwischen zwei möglichen Werten getroffen. Diese Werte sind um so eindeutiger, je größer der Unterschied zwischen ihnen ist (die Glühlampe brennt hell oder gar nicht).

In diesem Lehrbuch wollen wir die sogenannte **negative Logik** anwenden, da sie heute am gebräuchlichsten ist. Das bedeutet, daß am Ausgang des FF am leitenden Transistor das Signal 1 ($\hat{=}$ 0 V) erscheint und am gesperrten Transistor das Signal 0 ($\hat{=}$ +9 V). Da wir in unserem Beispiel nach Abb. 8.8 NPN-Transistoren verwenden, heißt dies:

Signal 1 entspricht 0 Volt und

Signal 0 entspricht +9 Volt.

Wir müssen uns diese Festlegung genau merken; sie soll nicht nur bei der Behandlung der verschiedenen FF, sondern auch für die Grundschaltungen der Logik im nächsten Abschnitt Gültigkeit haben.

Wenn wir ein FF in eine andere Lage bringen, es also **kippen** wollen, dann kann das auf zweierlei Art geschehen: Wir können entweder den gesperrten Transistor zum Leiten bringen oder aber den leitenden Transistor sperren. In unserem Beispiel haben wir zwei NPN-Transistoren verwendet. Wenn wir den zuerst gesperrten Transistor 2 zum Leiten bringen wollen, müssen wir seiner Basis positives Potential zuführen. Die Höhe des Potentials muß größer als die Diffusionsspannung der Basis-Emitter-Strecke sein; das sind etwa 0,7 Volt. Würde man die volle positive Spannung = +9 V anlegen (Signal 0), dann würde der Basisstrom so stark ansteigen, daß der Transistor zerstört würde. Wenn also ein gesperrter Transistor mit 0 angesteuert wird, muß zur Strombegrenzung in die Steuerleitung ein Schutzwiderstand eingebaut werden (vgl. hierzu auch Abb. 8.9).

Schaltungstechnisch ist es wesentlich einfacher, den vorher leitenden Transistor zu sperren. Das geschieht am einfachsten dadurch, daß an die Basis des leitenden Transistors Signal 1 angelegt wird, d.h., Basis und Emitter werden kurzgeschlossen. Da nach unserer Vereinbarung 0 Volt dem Signalwert 1 entspricht, heißt das, daß der leitende Transistor unseres FF mit Signal 1 gesperrt wird. Im Gegensatz zu vorher kann dabei der Transistor nicht überlastet werden. Schutzwiderstände in der Steuerleitung sind nicht notwendig. Wird das angelegte Signal wieder abgeschaltet, bleibt das FF in der eingenommenen Lage, bis ein weiteres Signal — jetzt aber an den anderen Eingang gelegt — das Zurückkippen einleitet.

Das Ansteuern der FF in der hier beschriebenen Weise, bei der mit Gleichspannung galvanisch das Steuersignal an die Basis eines Transistors gelegt wird, nennt man **statisches Ansteuern**, die Eingänge somit die **statischen Eingänge** eines FF.

Es gibt eine Reihe von Anwendungsgebieten, bei denen die FF nicht mehr galvanisch angesteuert werden können. Entweder verbieten es die Potentialverhältnisse, oder die Umschaltungen müssen mit sehr hoher Geschwindigkeit vollzogen werden. Hierfür hat man die sogenannten **dynamischen Eingänge** geschaffen, deren symbolische Darstellung Abb. 8.14 zeigt. In dynamischen Eingängen liegt eine Kapazität in Reihenschaltung.



Abb. 8.14 —
FF mit
dynamischen
Eingängen

Dynamische Eingänge geben das Steuersignal nur dann weiter, wenn es sich sprunghaft ändert. Dabei können die Eingänge wahlweise so geschaltet werden, daß sie bei einem Sprung von 1 nach 0 oder auch bei einem Sprung von 0 nach 1 das Umkippen des FF ermöglichen.

8.3.6. Anwendungsbeispiele für FF

8.3.6.1. Flipflop als Speicher

Wie wir schon eingangs erwähnt haben, bleibt das FF immer in der Lage, in die es vorher gesteuert worden ist. **Wenn also das FF einmal ein Steuersignal erhalten hat, dann speichert es den Schaltzustand dauernd.** Ein FF kann in der Ruhelage oder aber in der Arbeitslage sein. Der Ausgang A_2 hat z.B. den Zustand 1 oder 0. Mit anderen Worten: Da ein FF zwischen zwei Zuständen unterscheiden und diese Zustände auch speichern kann, wird es als Gedächtniszelle in digitalen Schaltverknüpfungen verwendet.

Der Begriff **digital** bedeutet in diesem Zusammenhang nichts anderes, als daß es Schaltkreise gibt, die nur in Stufen oder Sprüngen funktionieren, wie in unseren Überlegungen ein Ausgang eines FF nur die Unterscheidung zwischen 1 und 0 aufzeigen kann. Als Beispiel sind auch Digitaluhren zu nennen. Die Anzeige der Zeit springt bei einfachen Digitaluhren nur jede volle Minute weiter. Zwischenwerte sind nicht ablesbar.

Da ein FF nur den Informationsinhalt 1 oder 0 speichern kann, benötigt man für die Speicherung einer größeren Menge von Informationen — z.B. mehrstelligen Zahlen — eben mehrere FF, die nur sinnvoll aneinandergeschaltet werden müssen.

8.3.6.2. Flipflop als Frequenzteiler

Nachstehendes Impulsdigramm soll den Unterschied zwischen den Eingangs- und den Ausgangsimpulsen zeigen.

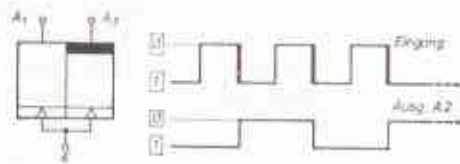


Abb. 8.15 — FF mit Ausgangs- und Eingangsimpulsen

Wir haben dem FF einen gemeinsamen dynamischen Eingang gegeben. Die Schaltung ist so ausgelegt, daß das FF bei einem Signalsprung von 0 nach 1 umkippt. Der von uns betrachtete Ausgang A_2 führt im Ruhezustand Signal 1. Mit dem ersten Sprung von 0 nach 1 am Eingang kippt das FF um, am Ausgang A_2 springt das Signal 1 nach 0. Wechselt das Eingangssignal von 1 nach 0, so bleibt das FF in der eingenommenen Lage. Erst mit dem nächsten Impuls, d.h. der nächste Sprung von 0 nach 1, wird das FF wieder in die Ruhelage zurückgesetzt. Jetzt erst erscheint am Ausgang A_2 wieder das Signal 1. Wir erkennen also deutlich, daß **nur bei jedem zweiten Eingangssignal das Ausgangspotential einmal wechselt**. Schalten wir nun zwei oder drei FF in der Art nach Abb. 8.16 hintereinander, so zeigt sich, daß FF II nur bei jedem zweiten Eingangsimpuls umkippt und FF III sogar nur bei jedem vierten Eingangsimpuls umkippt. Am Ausgang A_2 des FF III erscheinen die Signalzustände viermal so lang wie das Eingangssignal, oder — auf die Schaltfrequenz bezogen — die Ausgangsfrequenz der dreiteiligen FF-Kette beträgt nur $\frac{1}{4}$ der angelegten Steuerfrequenz.

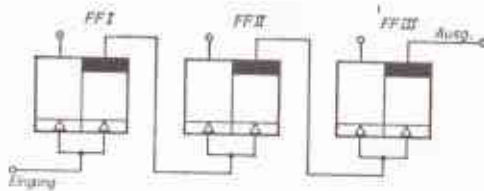


Abb. 8.16 — FF-Kette als Frequenzteiler

Je mehr FF in der beschriebenen Weise hintereinandergeschaltet werden, um so weiter kann eine angelegte Signalfrequenz geteilt werden. Von Stufe zu Stufe wird sie jeweils auf die Hälfte herabgesetzt.

8.3.6.3. Flipflop als Zähler

Wenn wir vier FF mit dynamischen Eingängen in der nach Abb. 8.17 gezeigten Weise hintereinanderschalten, dann wissen wir aus dem vorhergehenden Abschnitt, daß die Frequenz der Eingangsimpulse von Stufe zu Stufe jeweils durch 2 geteilt wird.

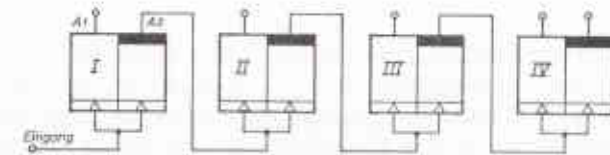


Abb. 8.17 — Zählschaltung mit 4 FF

Bei vier FF ist die Frequenz der letzten Stufe nur noch $\frac{1}{8}$ der Eingangsfrequenz.

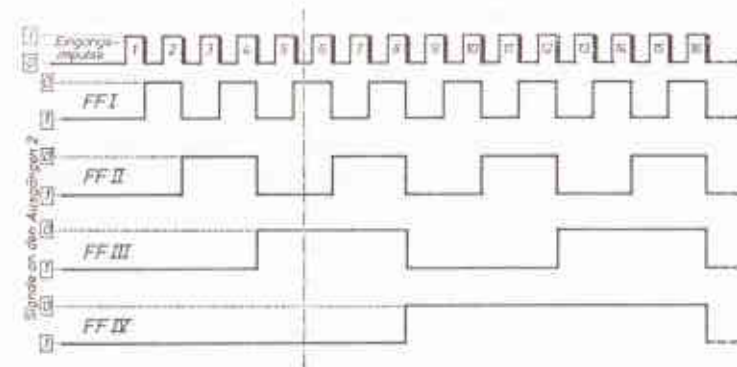


Abb. 8.18 — Impulsdiagramm der Zählschaltung

In Abb. 8.18 sind nun die Impulse aller FF untereinander gezeichnet. Aus den Impulsen kann aber nun die Stellung jedes einzelnen FF abgelesen werden. Wenn wir z.B. nach 5 Eingangsimpulsen aus dem Impulsdiagramm ablesen:

FF I	:	A_2	=	0
FF II	:	A_2	=	1
FF III	:	A_2	=	0
FF IV	:	A_2	=	1

so ist

FF I in Arbeitslage,
FF II in Ruhelage,
FF III in Arbeitslage und
FF IV in Ruhelage.

Sie können jetzt selbst einmal alle möglichen Schaltzustände der einzelnen FF untersuchen und werden sehr bald erkennen, daß bis (einschließlich) zum Eingangsimpuls Nr. 15 jeweils eine ganz bestimmte Stellung der FF zueinander erreicht wird, eine Stellung, die unverwechselbar und nur einmal vorkommt. Erst der 16. Impuls stellt alle FF in die Ruhelage zurück und entspricht damit der Nullage (Ausgangslage).

Mit vier solcher FF können wir daher bis 16 zählen, wenn wir unterstellen, daß die Nullage aller FF die Dezimalzahl 0 bedeutet. Nehmen wir ein weiteres FF dazu, können wir bis 32, mit 6 FF bis 64, mit 7 FF bis 128 usw. zählen. Wie wir bereits erwähnt haben, kann ein FF auch als Speicher verwendet werden. Es ist also möglich, eine einmal auf diese Weise abgezahlte Impulsfolge zu speichern.

In Datenverarbeitungsanlagen, Systemen der Steuer- und Regeltechnik, in der Elektrotechnik usw. sind immer sehr viele FF in der einen oder anderen Art enthalten.

8.4. Monostabile Kippstufe / Schmitt-Trigger

Unter einer **monostabilen Kippstufe** ist eine Kippstufe mit **einseitiger Ruhelage** zu verstehen. Diese Ruhelage verläßt sie nach einem von außen zugeführten Impuls, dessen Länge beliebig sein kann, und kehrt dann nach einer ganz bestimmten Zeit wieder ohne äußeres Zutun in die Ruhelage zurück. Der Zeitraum, in dem die Schaltung in der Arbeitslage verbleibt, wird durch die Zeit bestimmt, während der sich ein Kondensator über einen bestimmten Widerstand entladen kann. Für die Dauer der Arbeitslage der Kippstufe wird am Ausgang A der auf eine fest abgestimmte Länge gebrachte Impuls erscheinen. Der Ausgangsimpuls ist unabhängig von der Form des Eingangsimpulses immer ein Rechteckimpuls. Abb. 8.19 zeigt das Schaltsymbol und Impulsdiagramm der monostabilen Kippstufe.

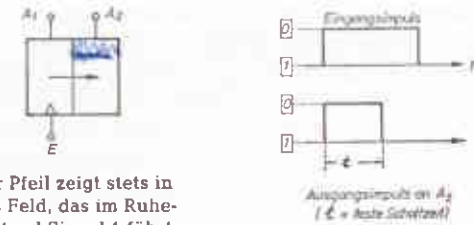


Abb. 8.19 — Symbol und Diagramm der monostabilen Kippstufe

Monostabile Kippstufen können zur Impulssteuerung verwendet werden. In ihrer Dauer verschieden lange Impulse werden durch die Schaltung auf feste Impulszeiten gebracht. Die Impulse werden durch die monostabile Kippstufe regeneriert. Eine speziell für die Impulserneuerung aufgebaute Kippstufe wird als **Schmitt-Trigger bezeichnet**. Wie bereits im Abschn. 4.6 „Der Transistor als Schalter“ festgestellt wurde, darf ein Schalttransistor nur die Zustände EIN oder AUS einnehmen. Für alle Zwischenstadien würden Schalttransistoren überlastet und somit zerstört. Da bei der Übertragung von Rechteck-Steuerimpulsen durch Kapazitäten und Widerstände der Schaltung Verformungen (Verzerrungen) der Rechteckimpulse verursacht werden, müssen sie zur Ansteuerung von Schalttransistoren regeneriert werden. Mit Hilfe eines Schmitt-Triggers gelingt es, verformte Steuerimpulse wieder in Rechteckimpulse mit steilen Anstiegs- und Abfallflanken zurückzuformen.

Monostabile Kippstufen sind weiterhin in der Lage, mehrere unerwünschte kurzzeitige Impulse zu unterdrücken, wie sie zum Beispiel bei Kontaktprellungen entstehen können. So können sie elektronischen, aus Flipflops bestehenden Zehlschaltungen vorgeschaltet werden, wenn von mechanischen Kontakten stammende Zählimpulse aufgenommen werden sollen.

Die Ausgangsspannung (Gleichspannung) eines Schmitt-Triggers ändert sich sprunghaft zwischen 0 V und der Speisespannung $+U$, wenn der Eingang mit zunehmender Spannung angesteuert wird. Den Schaltungsaufbau zeigt Abb. 8.20.

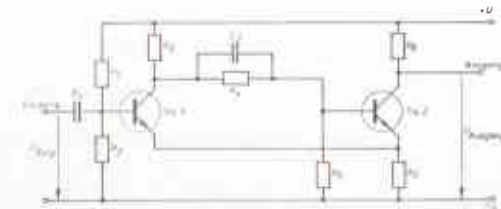


Abb. 8.20 — Prinzipschaltbild des Schmitt-Triggers

Die Schaltung besteht im Prinzip aus einem zweistufigen Transistorverstärker. Der Ausgang der 1. Stufe ist dabei über den Widerstand R_4 mit dem Eingang der 2. Stufe verbunden (Gleichstromkopplung). Die Basis-Emitter-Spannung des Transistors 1 wird durch den Spannungsteiler R_1/R_2 , die Basis-Emitter-Spannung des Transistors 2 über den Spannungsteiler R_4/R_5 gewonnen, der die Spannung am Kollektor des Transistors 1 abnimmt. Beide Transistoren haben den Emittewiderstand R_7 gemeinsam.

In der Regel ist die Schaltung so ausgelegt, daß im Ruhezustand der Transistor 1 gesperrt und der Transistor 2 leitend ist.

Der Spannungsteiler R_1/R_2 wird deshalb so bemessen, daß die Basis-Emitter-Spannung U_{BE1} noch unterhalb der Diffusionsspannung des PN-Übergangs Basis-Emitter liegt. Das ist bei Siliziumtransistoren etwa unterhalb 0,6 V (vgl. hierzu auch Abschn. 4.3.1). Dabei muß aber noch der Spannungsabfall am Widerstand R_7 berücksichtigt werden.

Somit wird $U_{BE1} = I_2 \cdot R_2 - I_7 \cdot R_7$.
Zur Verdeutlichung siehe Schaltungsauszug a (Abb. 8.21).



Abb. 8.21 — Schaltungsauszug a

Am Kollektor des gesperrten Transistors 1 liegt praktisch die volle Speisespannung $+U$. Diese wird durch den Spannungsteiler R_4/R_5 abgegriffen und der Basis des Transistors 2 zugeführt. Bei richtiger Dimensionierung der Bauteile erhält der Transistor 2 eine so hohe Basisspannung, daß er sicher leitet.

Am Ausgang des leitenden Transistors 2 liegt fast 0 V an (in Wirklichkeit die kleine Spannung $I_7 \cdot R_7$). Der Schaltungsauszug b (Abb. 8.22) zeigt die Zuführung der Basisspannung für den Trs 2.

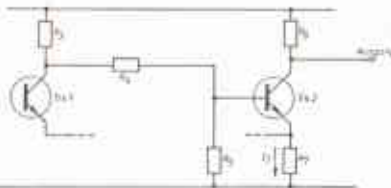


Abb. 8.22 — Schaltungsauszug b

Der Schaltzustand bleibt so lange bestehen, bis der Eingang mit Wechsellspannung angesteuert wird.

Wird an den Kondensator C_1 eine Wechsellspannung gelegt, so hebt die positive Halbwelle die Basisspannung des Transistors 1 an, bis U_{BE1} den Wert der Diffusionsspannung überschreitet. Jetzt wird der Transistor 1 leitend, wodurch seine Kollektorspannung auf nahezu 0 V abfällt. Damit sinkt die Basis-Emitter-Spannung des Transistors 2 schlagartig, so daß er sperrt. Am Kollektor des Transistors 2 — und damit am Ausgang — liegt jetzt die volle Speisespannung U .

Nimmt die Eingangsspannung wieder ab, sinkt auch die Basis-Emitter-Spannung des Transistors 1, bis er wieder sperrt. Am Kollektor des Trs 1 erscheint die Spannung $+U$. Da die Basis-Emitter-Spannung des Transistors 2 am Kollektor des Transistors 1 abgegriffen wird, geht der Transistor 2 in den leitenden Zustand über. Am Ausgang liegt wieder fast 0 V.

Es ist wichtig zu wissen, daß die Schaltvorgänge sehr schnell vor sich gehen, so daß die Ausgangsspannung zwischen den beiden möglichen Werten 0 und $+U$ sprunghaft wechselt. Legt man an den Eingang eine Sinusspannung, so können am Ausgang Rechteckimpulse abgenommen werden (vgl. hierzu Abb. 8.23).

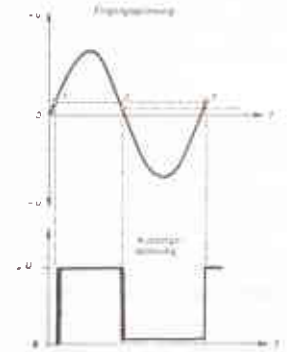


Abb. 8.23 — Schmitt-Trigger als Impulsformer:
Eingangs- und Ausgangsspannung

Aus der Abb. 8.23 ist weiterhin zu sehen, daß der Schmitt-Trigger erst von einer gewissen Schwellspannung ab umschaltet. Die Höhe dieses Wertes kann durch geeignete Einstellung der Basisvorspannung U_{BE1} (Spannungsteilerverhältnis R_1/R_2 unter Beachtung des Spannungsabfalls an R_7) vorbestimmt werden. Damit wird erreicht, daß z.B. unterhalb der Schwellspannung liegende Störimpulse die Schaltung nicht ansprechen lassen; sie wird dadurch „störsicher“. Weiterhin ist zu sehen, daß der Trs 1 etwas unterhalb des Einschaltwertes ausschaltet.

Läßt man den Eingangskondensator C_1 weg, so können auch Gleichspannungen oder sich sehr langsam ändernde Spannungen die Schaltung zum Kippen bringen. In dieser Ausführungsform arbeitet der Schmitt-Trigger als **Schwellwertschalter**. Als Beispiel sehen Sie in Abb. 8.24 eine langsam ansteigende Eingangsspannung, die bei einem vorgewählten Schwellwert den Schmitt-Trigger umschaltet.

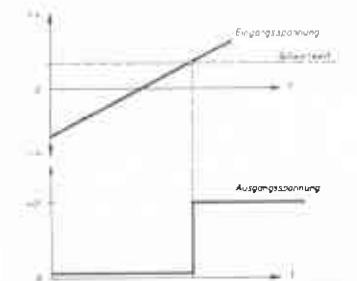


Abb. 8.24 — Schmitt-Trigger als Schwellwertschalter:
Eingangs- und Ausgangsspannung

8.5. Astabile Kippstufe

In Abb. 8.25 ist die Schaltung einer astabilen Kippstufe dargestellt.

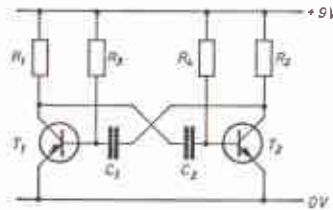


Abb. 8.25 — Astabile Kippstufe

Eine astabile Kippstufe hat keine stabile Lage; sie ändert ihren Schaltzustand laufend selbständig. An ihren Ausgängen können Rechteckimpulse abgenommen werden, deren Frequenz durch die Größe der Bauteile R_3 , C_1 und R_4 , C_2 festgelegt ist (Zeitkonstante).

Die Schaltung besteht im wesentlichen aus zwei Transistoren, die so miteinander gekoppelt sind, daß die Basis jedes Transistors mit dem Kollektor des anderen Transistors verbunden ist. Schaltet ein Transistor (Trs 1) durch, wird er also leitend, dann gelangt Emittential (also 0 Volt) an die linke Seite des Kondensators C_2 . Im Einschaltaugenblick wirkt ein Kondensator wie ein Kurzschluß (0 Ohm). Somit liegt zumindest kurzzeitig auch Emittential (ca. 0 Volt) an der Basis des Trs 2. Dadurch wird dieser sofort gesperrt. Der Vorgang dauert aber nur eine bestimmte Zeit, nämlich so lange, bis der Kondensator C_2 über den Widerstand R_4 aufgeladen worden ist. Jetzt kommt Trs 2 wieder zum Leiten, wodurch Trs 1 auf die gleiche Weise, wie vorher beschrieben, über den Kondensator C_1 gesperrt wird. Dieser Vorgang wiederholt sich laufend, so daß eine astabile Kippstufe wie ein Generator dauernd Rechteckimpulse abgibt. Es ergibt sich deutlich, daß die Dauer eines Impulses von zwei Größen abhängt:

- von der Kapazität der Kopplungskondensatoren C_1 und C_2 und
- von der Größe der Widerstände R_3 und R_4 (Zeitkonstante).

Je größer R_3 und R_4 sowie C_1 und C_2 , um so länger dauern die Kippvorgänge, d.h., um so kleiner ist die Impulsfrequenz.

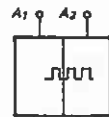


Abb. 8.26 zeigt die symbolisierte Darstellung einer astabilen Kippstufe.

Abb. 8.26

8.6. Selbstbau einer Kippstufe

Um die Kippstufe in ihrer Wirkungsweise beobachten zu können, wollen wir einen Blinkgeber aufbauen, der aus einer astabilen Kippstufe (im Prinzip in Abb. 8.25 gezeigt) und einer Transistorschaltstufe besteht. Der Schalttransistor hat dann die Aufgabe, die durch die Kippstufe erzeugten Impulse so weit zu verstärken, daß eine Glühlampe im Rhythmus der Impulsfrequenz aufleuchten kann.

Zu Beginn werden die Transistortypen und die Speisespannung U_B festgelegt und als Transistoren die uns schon bekannten NPN-Typen BC 107 verwendet. Von der bisher üblichen Speisespannung 9 V müssen wir allerdings abgehen, da Glühlampen für 9 V im Handel nicht erhältlich sind. Wir entscheiden uns deshalb für eine Betriebsspannung von 6 V. Sie ergibt sich z.B. durch Reihenschaltung von vier Monozellen mit je 1,5 V Spannung.

Zur Berechnung der einzelnen Schaltelemente gehen wir von den Daten des Transistors BC 107 aus; die wichtigsten sind:

$$\text{maximal zulässiger Kollektorstrom } I_{C_{\max}} = 100 \text{ mA,}$$

$$\text{Stromverstärkungsfaktor (Mindestwert) } B = 125.$$

Da die Kippschaltung symmetrisch ist, sind die für eine Stufe ermittelten Werte auf die zweite ohne Änderung zu übernehmen. Der Kollektorstrom eines Transistors im leitenden Zustand soll im Hinblick auf eine spätere Ausgangsbelastung möglichst groß sein, jedoch den Grenzwert $I_{C_{\max}}$ mit Sicherheit nicht überschreiten. Wir nehmen zunächst

$$I_C = 20 \text{ mA}$$

an, denn dieser Wert wird beiden Forderungen in etwa gerecht.

Die Größe des Arbeitswiderstands R_2 ergibt sich aus der Formel:

$$R_2 = \frac{U}{I_C} = \frac{6 \text{ V}}{20 \text{ mA}} = \frac{6 \text{ V}}{0,02 \text{ A}} = 300 \Omega.$$

Wir wählen aus der Widerstandsnormreihe den Wert $R_2 = 330 \Omega$.

Der tatsächliche Kollektorstrom I_C ergibt sich dann zu:

$$I_C = \frac{U}{R_A} = \frac{6 \text{ V}}{330 \Omega} = 0,018 \text{ A}$$

$$I_C = 18 \text{ mA.}$$

Aus R_2 kann der Basisvorwiderstand R_4 bestimmt werden, der den Transistor im völlig leitenden Zustand hält; es gilt die Näherungsformel für Si-Schalttransistoren:

$$R_4 \approx 0,8 R_2 \cdot B.$$

Die Formel ist aus folgender Überlegung abgeleitet: Bei völlig leitenden Transistoren beträgt der Spannungsabfall des Kollektorstroms am Arbeitswiderstand $R_2 \cdot I_C$ und der des Basisstroms am Basisvorwiderstand $R_4 \cdot I_B$ (vgl. hierzu Abb. 8.27).

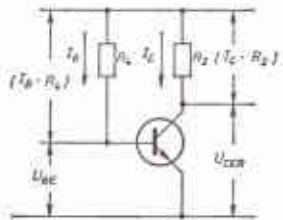


Abb. 8.27 — Blinkgeber (Schaltungsausschnitt 1)

Für die Transistorschaltung 8.27 gilt ferner:

$$R_4 \cdot I_B = U - U_{BE}$$

$$R_2 \cdot I_C = U - U_{CE\text{Rest}}$$

Da U_{BE} und $U_{CE\text{Rest}}$ etwa gleich groß sind, gilt

$$R_4 \cdot I_B \approx R_2 \cdot I_C$$

Die Formel wird nach R_4 umgestellt und lautet:

$$R_4 = R_2 \cdot \frac{I_C}{I_B}$$

$$\frac{I_C}{I_B} = B$$

$$R_4 = R_2 \cdot B$$

Bei der so errechneten Größe von R_4 wird der Transistor gerade leitend. Mit Sicherheit leitend ist er nur dann, wenn R_4 etwas kleiner gewählt wird; das ergibt sich aus der oben bereits angegebenen Formel:

$$R_4 = 0,8 \cdot R_2 \cdot B$$

R_2 wurde errechnet und festgelegt auf $330 \, \Omega$, und B beträgt bei dem Transistor BC 107 mindestens 125.

$$R_4 = 0,8 \cdot 330 \, \Omega \cdot 125 = 33\,000 \, \Omega = 33 \, \text{k}\Omega$$

Damit liegt auch die Größe des Basisvorwiderstands fest. Durch die Koppelkondensatoren wird die Kippfrequenz der Schaltung bestimmt. Als Blinkfrequenz ist der Wert von etwa 2 Hz recht günstig. 2 Hz bedeutet, daß pro Sekunde eine Lampe zweimal aufleuchtet (Abb. 8.28).

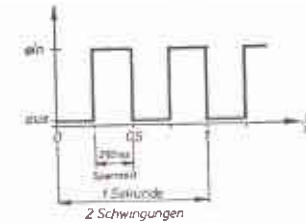


Abb. 8.28 — Blinkrhythmus

Aus Abb. 8.28 geht hervor, daß die Sperrphase t eines Transistors 250 ms betragen muß. Für die Ermittlung der Dauer der Sperrphase gilt die Erfahrungsformel

$$t = 0,7 \cdot R_4 \cdot C_K$$

Darin bedeuten R_4 = Basisvorwiderstand und

C_K = Kapazität des Koppelkondensators in F.

Stellen wir die Formel nach C_K um, so kann die Kapazität direkt aus t und R_4 errechnet werden.

$$C_K = \frac{t}{0,7 \cdot R_4} = \frac{250 \cdot 10^{-3} \text{ s}}{0,7 \cdot 33 \cdot 10^3 \frac{\text{V}}{\text{A}}} = 10,8 \cdot 10^{-6} \text{ F} = 10,8 \, \mu\text{F}$$

Wir wählen für C_K einen gängigen Wert: $C_K = 10 \, \mu\text{F}$.

Damit wird zwar die Kippfrequenz geringfügig größer, das ist aber bedeutungslos. Als Kondensator eignet sich am besten ein Elektrolytkondensator, da seine Baugröße sehr gering ist.

Zum Schluß ist nun noch der Lampenschaltverstärker zu berechnen. Seine Dimensionierung hängt in erster Linie von der zu schaltenden Lampenleistung ab. An brauchbaren Glühlampen bieten sich an:

6 V/0,1 A mit Schraubgewinde E 10 oder

6 V/0,05 A mit Schraubgewinde E 10.

Wir sehen für unsere Blinkschaltung die stärkere der beiden Lampen vor, also 6 V/0,1 A. So bleibt die Möglichkeit offen, auch zwei Lampen gleichzeitig zu schalten; es müßte dann nur die Ausführung mit 0,05 A gewählt werden; Abb. 8.29 zeigt den Lampenverstärker.

Im eingeschalteten Zustand fließt bei einer Speisespannung U_S von 6 V ein Lampenstrom $I_L = 0,1 \text{ A}$. Dieser bildet den Kollektorstrom des Transistors, also ist $I_C = 0,1 \text{ A}$. Als maximaler Kollektorstrom ist 0,1 A für den Transistor BC 107 gerade noch möglich, so daß wir

auch hier diesen Typ einsetzen können. Der Vorwiderstand R_V in der Basisleitung setzt die am Ausgang der Kippstufe stehende Spannung von 6 V (während der Sperrphase des Transistors) auf die erforderliche Basis-Emitter-Spannung U_{BE} herab. Damit der Schalttransistor sicher leitet, nehmen wir U_{BE} mit 1 V an, so daß der Spannungsabfall U_R an R_V durch den Basisstrom I_B 5 V betragen muß. Der Basisstrom errechnet sich aus der umgestellten Formel für die Stromverstärkung B .

$$I_B = \frac{I_C}{B} = \frac{0,1 \text{ A}}{125} = 0,0008 \text{ A} = 0,8 \text{ mA}.$$

Dieser den Ausgang der Kippstufe belastende Strom von 0,8 mA ist gegenüber dem in der Kippschaltung fließenden Kollektorstrom von etwa 18 mA sehr klein und beeinflusst so die Kippstufe nicht.

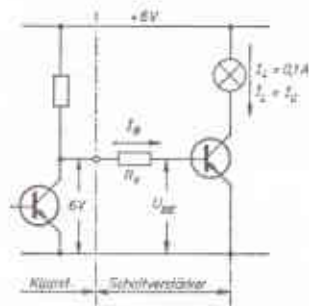


Abb. 8.29 — Blinkgeber (Schaltungsausschnitt 2)

Der Vorwiderstand R_V wird so bemessen, daß $I_B = 0,8 \text{ mA}$ an ihm einen Spannungsabfall $U_R = 5 \text{ V}$ erzeugt.

$$R_V = \frac{U_R}{I_B} = \frac{5 \text{ V}}{0,0008 \text{ A}} = 6250 \Omega = 6,25 \text{ k}\Omega.$$

Mit Rücksicht auf ein sicheres Durchschalten des Schalttransistors wählen wir aus der Normreihe einen kleineren Wert:

$$R_V = 5,1 \text{ k}\Omega.$$

Mit diesem letzten Wert ist die Schaltung des elektronischen Blinkgebers mit allen Bauelementen dimensioniert und in Abb. 8.30 dargestellt.

Da der Kollektorstrom des Transistors BC 107 an der oberen zulässigen Grenze liegt, empfiehlt es sich, den Schalttransistor T_3 durch ein Kühlblech zu kühlen.

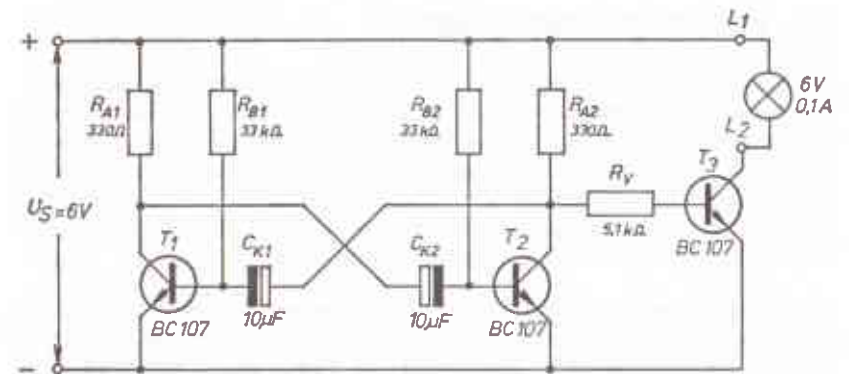


Abb. 8.30 — Blinkgeber (dimensionierte Schaltung)

In der dargestellten Schaltung sind die Widerstände anders bezeichnet, als in der Berechnung. Es gilt

$$R_{B1} \text{ bzw. } R_{B2} = R_4 \text{ und} \\ R_{A1} \text{ bzw. } R_{A2} = R_2.$$

Wir wollen die Schaltung wieder auf einer Pertinaxplatte aufbauen. Sie findet auf einer Platine mit den Abmessungen 70 x 35 mm Platz. Wie die Bauteile befestigt und untereinander verdrahtet werden, geht aus den Abb. 8.31 und 8.32 hervor.

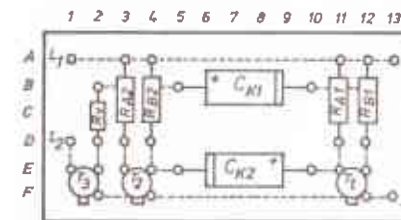


Abb. 8.31 — Platine (Bauteilseite)

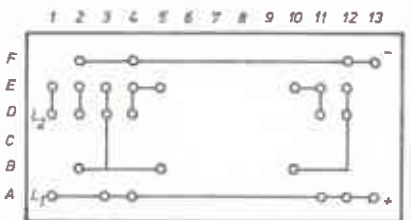


Abb. 8.32 — Platine (Verdrahtungsseite)

Für den Aufbau des Blinkgebers benötigte Teile:

- 3 Transistoren BC 107
- 2 Widerstände 330 Ω 0,5 W
- 2 Widerstände 33 k Ω 0,25 W
- 1 Widerstand 5,1 k Ω 0,25 W
- 2 Elektrolytkondensatoren 10 μF 12/15 V
- 1 Glühlampe 6 V 0,1 A Schraubgewinde E 10
- 1 Schraubfassung E 10
- 1 Platine (kupferbeschichtet) 70 x 35 mm mit 5-mm-Lochraster

9. Verknüpfungsglieder

Verknüpfungsglieder, die auch als logische Schaltungen oder Gatter bezeichnet werden, dienen in der Elektronik zur Datenverarbeitung. Sie ordnen die vielen eingegebenen Signale nach ganz bestimmten, festgelegten mathematischen Regeln einander zu. Die logischen Schaltungen stellen also gewissermaßen für alle eingehenden Signale die Weichen und ermöglichen bestimmte Rechenoperationen, wie Addition, Subtraktion, Multiplikation usw.

Die Schaltungen der digitalen Technik werden heute allgemein mit elektronischen Bauelementen ausgeführt; das sind Widerstände, Dioden, Transistoren und Kondensatoren. Schaltungen mit sämtlichen Einzelbauelementen werden aber zu unübersichtlich; daher werden Baugruppen ähnlich wie FF durch vereinfachte Schaltsymbole ersetzt.

Wird daher eine Schaltung entworfen, so werden die Schaltaufgaben und Schaltbedingungen zunächst mit Hilfe der allgemeinen Schaltsymbole gezeichnet. Die Schaltsymbole sind genormt; alle möglichen Schaltvarianten sind grundsätzlich auf drei Grundschaltungen zurückzuführen.

9.1. UND-Glied (UND-Gatter)

Abb. 9.1 zeigt das Symbol eines UND-Glieds mit zwei Eingängen.



Abb. 9.1 — UND mit zwei Eingängen

Am einfachsten ist die Schaltung mit Relais zu verwirklichen.

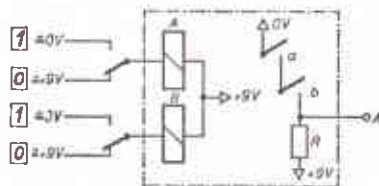


Abb. 9.2 — UND mit Relais

Wir ersehen aus der Abb. 9.2, daß im Ruhezustand am Ausgang der Schaltung +9 Volt liegt. Wenn nun Relais A oder Relais B allein anspricht, ändert sich am Ausgangspotential nichts. **Erst wenn die Relais A und B gleichzeitig ansprechen, also die Kontakte a und b gleichzeitig geschlossen sind, liegt am Ausgang 0 Volt**, d.h., am Ausgang liegt nach unserer Vereinbarung **Signal 1**. Wir können dieses Beispiel auf unser Schaltsymbol des UND-Glieds (auch UND-Gatter genannt) übertragen und können sagen: Wenn an beide Eingänge Signal 1 angelegt wird, dann erscheint am Ausgang ebenfalls Signal 1. Diese Aussage trifft auch zu, wenn ein UND-Glied mehrere Eingänge besitzt.

Das UND-Glied kann mit Dioden aufgebaut werden.

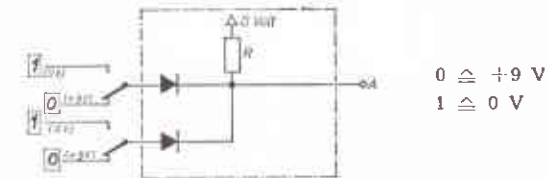


Abb. 9.3 — UND mit Dioden

Wenn beide Eingänge Signal 0 führen (das entspricht +9 V), erscheint am Ausgang +9 V, das entspricht wieder dem Signalwert 0. Wenn einer der beiden Eingänge allein Signal 1 zeigt, wird die positive Spannung des anderen Eingangs immer noch zum Ausgang A durchgreifen, d.h., am Ausgang liegt immer noch Signal 0. Erst wenn beide Eingänge zugleich Signal 1 führen, verschwindet das positive Potential und dafür wirkt sich das Nullpotential über den Widerstand R auf den Ausgang aus; es liegt Signal 1 am Ausgang.

Finden Transistoren Verwendung, so kommt man zu nachstehender Schaltung:

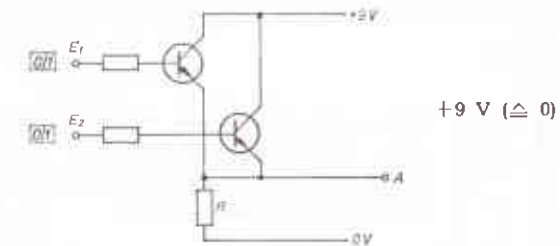


Abb. 9.4 — UND-Glied mit Transistoren

Wie zu erkennen ist, sind am Eingang die mechanischen Kontakte weggelassen worden. Damit wir die Wirkung der Schaltung besser übersehen können, nehmen wir die nachstehende Tabelle zu Hilfe, die Funktionstabelle genannt wird. Dabei wollen wir festhalten, daß an den Eingängen nur die beiden Signalzustände 0 oder 1, also keine Zwischenwerte oder gar ein spannungsloser Zustand, auftreten können.

Funktionstabelle

UND-Glied

E_1	E_2	A
0	0	0
1	0	0
0	1	0
1	1	1

Die Tabelle sagt folgendes aus:

Am Ausgang A erscheint nur dann Signal 1, wenn gleichzeitig der Eingang E_1 und der Eingang E_2 Signal 1 führen.

Diese Erkenntnis können wir auf alle UND-Glieder mit beliebig vielen Eingängen ausdehnen. **Wir wollen uns merken: Am Ausgang eines UND-Glieds liegt nur dann Signal 1, wenn alle Eingänge gleichzeitig Signal 1 führen.**

9.2. ODER-Glied (ODER-Gatter)

Diese Schaltung ist die zweite der Verknüpfungsglieder; ihr Symbol ist in Abb. 9.5 dargestellt; der Einfachheit halber wieder mit nur zwei Eingängen.



Abb. 9.5 — ODER-Glied mit zwei Eingängen

Da das Prinzip einer Funktionstabelle für ein UND-Glied aus dem vorhergehenden Abschnitt bekannt ist, kann die entsprechende Tabelle für das ODER-Glied bereits abgeleitet werden:

Funktionstabelle

ODER-Glied

E_1	E_2	A
0	0	0
1	0	1
0	1	1
1	1	1

Wir sehen aus dieser Tabelle, daß der Ausgang immer dann das Signal 1 führt, wenn einer der beiden Eingänge allein oder beide zusammen Signal 1 führen.

Daraus läßt sich wieder die Regel für alle ODER-Glieder ableiten. **Wir wollen uns merken: Am Ausgang eines ODER-Glieds liegt dann Signal 1, wenn mindestens ein Eingang Signal 1 führt.** Wird die Schaltung mit Relais aufgebaut, so ergibt sich ein Bild nach Abb. 9.6.

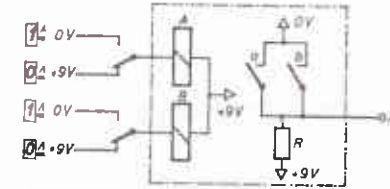


Abb. 9.6 — ODER-Glied mit Relais

Selbstverständlich können ODER-Glieder auch mit Dioden hergestellt werden (Abb. 9.7).

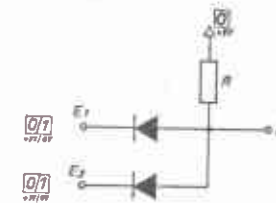


Abb. 9.7 — ODER-Glied mit Dioden

Wenn an den Eingängen Signal 0 liegt, wirkt die positive Spannung von +9 Volt, also ebenfalls Signal 0 am Ausgang. Wechselt an einem Eingang das Potential, erscheint also an E_1 oder an E_2 das Signal 1, dann greift das Signal zum Ausgang durch.

Als Beispiel ist in Abb. 9.8 noch ein ODER mit NPN-Transistoren dargestellt.

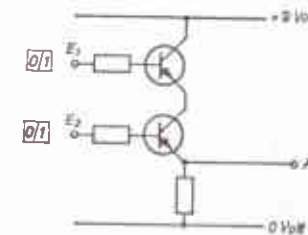


Abb. 9.8 — ODER-Glied mit Transistoren

Liegt an beiden Eingängen Signal 0, dann sind beide NPN-Transistoren durch das positive Potential an ihren Basen leitend und +9 Volt ($\hat{=}$ Signal 0) liegt am Ausgang. Wird an eine Basis Signal 1 angelegt, dann sperrt der betreffende Transistor und das Nullpotential ($\hat{=}$ Signal 1) wird über Widerstand R am Ausgang wirksam.

9.3. NICHT-Glied (NICHT-Gatter)

Ein NICHT-Glied hat die Aufgabe, ein anliegendes Signal in sein Gegenteil umzuwandeln. Es wird aus Signal 1 dann Signal 0 und umgekehrt. Wird die Umkehrschaltung mit Halbleiterbauteilen und nicht mit Relais ausgeführt, benötigt man dazu Transistoren; damit kann gleichzeitig eine Signalverstärkung erzielt werden. Die Umkehrschaltung nur mit Widerständen und Dioden aufzubauen, ist nicht möglich.

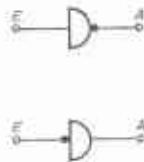


Abb. 9.9 — Übliche Symbole für das NICHT-Glied

Wird das NICHT-Glied mit Hilfe eines Relais hergestellt, so ergibt sich folgendes Bild:

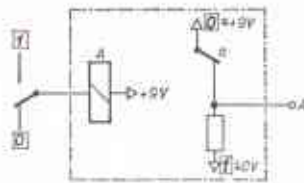


Abb. 9.10 — NICHT-Glied mit Relais

Wie Sie sehen, legt der Ruhekontakt a so lange Signal 0 an den Ausgang, bis das Relais A anspricht, den a-Kontakt öffnet und über den Widerstand das Signal 1 am Ausgang erscheint.

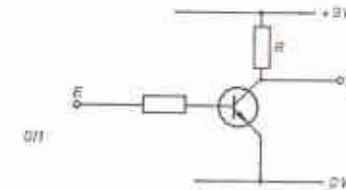


Abb. 9.11 — NICHT-Glied mit Transistor

Das NICHT-Glied mit einem NPN-Transistor nach Abb. 9.11 arbeitet wie folgt: Liegt am Eingang eine Spannung von +9 Volt ($\hat{=}$ Signal 0), dann ist der Transistor leitend und schaltet 0 Volt ($\hat{=}$ Signal 1) an den Ausgang. Wird die Basis mit Signal 1 gesperst, dann unterbricht der Transistor die Verbindung zu 0 Volt, indem er sehr hochohmig wird, und das Signal 0 (+9 V) erscheint am Ausgang.

Eine normale Verstärkerstufe mit einem Transistor in Emitterschaltung kehrt also immer das Signal um. In der linearen Technik — der üblichen Verstärkertechnik — ist dies als selbstverständlich bekannt; man spricht dort von einer Phasenverschiebung zwischen Ein- und Ausgang einer Verstärkerstufe um 180° .

Wie eingangs erwähnt, bauen sich alle möglichen Verknüpfungen aus den besprochenen drei Grundschaltungen auf. Selbstverständlich werden die Schaltungen anders aussehen, wenn sie anstatt mit NPN- mit PNP-Transistoren ausgeführt werden oder wenn anstelle der negativen Logik positive Logik vereinbart wird. Man darf also nicht überrascht sein, wenn in industriellen Schaltungen — besonders bei älteren Ausgaben — Abweichungen von unseren Schaltbeispielen zu finden sind.

10. Anhang

10.1. Die Bezeichnung von Halbleiter-Bauelementen

Jeder Hersteller von Halbleiter-Bauelementen führte zunächst in der stürmischen Entwicklung der Halbleitertechnik eine betriebseigene Bezeichnung ein. Inzwischen hat sich jedoch eine einheitliche Bezeichnung nach vorwiegend drei Bezeichnungsschlüsseln durchgesetzt.

In Deutschland werden Halbleiter-Bauelemente nach folgendem Schlüssel gekennzeichnet:

1. Buchstabe (Kennzeichen für das Halbleitermaterial):

- A** Ausgangsmaterial Germanium
- B** Ausgangsmaterial Silizium

2. Buchstabe (Kennzeichen für Art und Verwendungszweck des Bauelements):

- A** Diode (auch Kapazitätsdiode)
- C** Niederfrequenz-Transistor (Tonfrequenzbereich)
- D** Niederfrequenz-Leistungstransistor
- E** Tunneldiode (Esaki-Diode)
- F** Hochfrequenz-Transistor
- H** Hallsonde
- K** Hallgenerator
- L** Hochfrequenz-Leistungstransistor
- P** lichtempfindliches Bauelement (z.B. Fotoelement)
- R** steuerbarer Gleichrichter (Thyristor)
- S** Schalttransistor
- T** steuerbarer Leistungs-Gleichrichter (Leistungsthyristor)
- U** Leistungs-Schalttransistor
- Y** Leistungs-Diode
- Z** Referenzdiode (auch Zenerdiode)

Für Industrie-Typen wird in der Bezeichnung noch ein **dritter Buchstabe X, Y oder Z** verwendet. Die den Buchstaben folgende Zifferngruppe stellt eine Art Entwicklungs-Nr. dar, ohne Rückschlußmöglichkeit auf die Eigenschaften eines Halbleiter-Bauelements — mit einer Ausnahme:

Bei **Z-Dioden** folgt auf die Serien-Nr., durch einen Bruchstrich getrennt, eine Bezeichnung wie z.B. /C 5 V 6. Sie gibt die Zenerspannung der Z-Diode an. C 5 V 6 bedeutet, daß die Zenerspannung dieser Diode 5,6 Volt beträgt. Um Verwechslungen zu vermeiden, wird anstelle des Kommas ein V gesetzt. Genauso bedeuten C 24 : $U_z = 24$ Volt oder D 8 V 2 : $U_z = 8,2$ Volt. Die Buchstaben C bzw. D vor dem Spannungsschlüssel kennzeichnen die Toleranz (C = 5 %, D = 10 %).

Beispiele:

AA 118	Germanium (A) — Diode (A)
AF 126	Germanium (A) — HF-Transistor (F)
BSX 74	Silizium (B) — Schalttransistor (S) für industrielle Anwendung (X)
BZX 85/C6V5	Silizium (B) — Z-Diode (Z) für industrielle Anwendung (X) Zenerspannung $U_z = 6,5$ V bei einer Toleranz von ± 5 % (C).

Neben dem angegebenen Schlüssel wird auch der **USA-Bezeichnungsschlüssel** häufig angewendet. Er besteht aus einer Ziffer 1 ... 4, dem großen Buchstaben N und einer nachfolgenden Zifferngruppe. Die Unterscheidung ist nicht so fein möglich, wie nach dem deutschen Schlüssel. Es bedeuten:

- 1 N ...** Dioden
- 2 N ...** Transistoren, Thyristoren
- 3 N ...** Tetroden, Binistoren und ähnl.
- 4 N ...** Vierschicht-Dioden

Nachfolgende Ziffern sind laufende Typen- bzw. Entwicklungs-Nr. Manchmal sind sie auch Kennzeichen der elektrischen Daten (je nach Hersteller).

Halbleiter-Bauelemente werden in **Japan** nach folgendem Schlüssel gekennzeichnet:

- 1 S ...** Dioden
- 2 S ...** Transistoren, Thyristoren;

darauf folgt ein Buchstabe und eine Zifferngruppe

- A** = HF-PNP-Transistor
- B** = NF-PNP-Transistor
- C** = HF-NPN-Transistor
- D** = NF-NPN-Transistor
- S** = PNP-Vierschicht-Halbleiter-Bauelemente

Die Zifferngruppe dient der Typen-Kennzeichnung.

Beispiel:

2 SC 281 = NPN-Hochfrequenztransistor.

Kleingleichrichter für die Netzstromversorgung von Geräten werden nach folgendem Schlüssel gekennzeichnet:

1. Buchstabe zur Kennzeichnung der **Schaltungsart**

- E** = Einwegschialtung
B = Brückenschaltung (Graetzschaltung)
M = Mittelpunktsschaltung (Zweiwegschialtung)
V = Spannungsverdopplungsschialtung (Villardschialtung)

2. Zahlenangabe zur Kennzeichnung der zulässigen **Nennanschlußspannung** in Volt; im allgemeinen der Effektivwert der zugeführten Wechselspannung.

3. Buchstabe zur Kennzeichnung der **Belastungsart**

- C** = Kondensatorbelastung (Siebkette)
R = Widerstandsbelastung

4. Zahlenangabe zur Kennzeichnung des zulässigen **Nenngleichstroms** in Milliampere (arithmetischer Mittelwert).

Beispiel:

B 30 C 500 = Brückengleichrichter,
 Anschlußspannung 30 V_{eff},
 Kondensatorbelastung,
 Nenngleichstrom 500 mA.

10.2. Die Bezeichnung von Elektronenröhren

Die verschiedenartigen Röhrentypen werden nach bestimmten Schlüsseln bezeichnet. Der Fachmann erkennt anhand der Bezeichnung Art und Verwendungszweck einer Röhre. Der in Deutschland für Elektronenröhren allgemeiner Anwendungsbereiche (z.B. Rundfunk- und Fernsehempfänger, Phonogeräte usw.) übliche Bezeichnungsschlüssel ist in folgender Übersicht zusammengestellt.

Die Röhrenbezeichnung setzt sich aus zwei oder mehr Buchstaben und einer angehängten Zifferngruppe zusammen. Der 1. Buchstabe gibt

Aufschluß über die Art der Röhrenheizung, der 2. und weitere Buchstaben kennzeichnen das Röhrensystem. Die Zifferngruppe stellt eine Art Entwicklungs-Nr. dar. Sie läßt aber bei modernen Röhren der E- und P-Serie Rückschlüsse auf die Art des Röhrensockels zu.

An 1. Stelle		An 2., 3. und weiterer Stelle	
	Heizart		Röhrensystem Anwendungsbeispiele
A	4 V $\sim$	A	Diode HF-Gleichrichter
B	180 mA = — —	B	Doppeldiode HF-Gleichrichter
C	200 mA $\cong$	C	HF-/NF-Triode Vorverstärker
D	1,2 . . . 1,4 V Batt. — —	D	NF-Leistungstriode Endverstärker
E	6,3 V $\cong$ — —	F	HF-/NF-Pentode (Fünfpolröhre) Vorverstärker
K	2 V Batt.	H	Hexode (Sechspolröhre) Mischröhre
P	300 mA $\cong$ — —	L	Leistungspentode Endverstärker
U	100 mA $\cong$ — —	M	Anzeigeröhre (z.B. Mag.Band) Abstim- oder Aussteuerungsanzeige
V	50 mA $\cong$ — —	Y	Leistungsdioden Netz- oder Hochspannungsgleichrichter
		Z	Leistungs-Doppeldiode Netz-Zweiweggleichrichter

Zeichenerklärung:

- $\sim$ Wechselstrom
 = Gleichstrom
 $\cong$ Allstrom
 Batt. Batteriebetrieb
- || Parallelschialtung der Heizfäden
 — — Reihenschialtung der Heizfäden
 HF Hochfrequenz
 NF Niederfrequenz (Tonfrequenz)

Die erste Ziffer gibt Aufschluß über den Röhrensockel:

- 9 Pico 7-Sockel (7 Stifte)
 4 Pico 8-Sockel (8 Stifte)
 8 Pico 9-Sockel (9 Stifte)
 2 Pico 10-Sockel (10 Stifte).

Beispiel:

EABC 80 = Röhre mit 6,3 V Heizspannungsbedarf (E)
und einer Diode (A),
einer Doppeldiode (B)
und einer Triode (C).
Der Röhrensockel hat 9 Anschlußstifte (80).

10.3. Farbcode für Widerstände und Kondensatoren

Der Widerstands- oder Kapazitätsnennwert wird heute nahezu ausschließlich durch einen Farbcode gekennzeichnet. Dabei ergeben die ersten beiden Farbringe oder -punkte die Ziffernfolge, der dritte den Erweiterungsfaktor und der vierte die Toleranz.

	Nennwert in Ω bzw. pF			Toleranz
Farbe	1. Ring (1. Ziffer)	2. Ring (2. Ziffer)	3. Ring (Faktor)	4. Ring
schwarz	—	0	1	—
braun	1	1	10^1	1 %
rot	2	2	10^2	2 %
orange	3	3	10^3	—
gelb	4	4	10^4	—
grün	5	5	10^5	—
blau	6	6	10^6	—
violett	7	7	10^7	—
grau	8	8	10^8	—
weiß	9	9	10^9	—
gold	—	—	10^{-1}	5 %
silber	—	—	10^{-2}	10 %
keine	—	—	—	20 %

Ist kein Toleranz-Kennring vorhanden, so wird jeweils von dem Farbring aus angefangen, der einem Ende des Widerstandskörpers am nächsten liegt.

Beispiel 1:

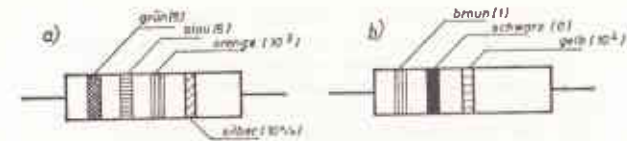


Abb. 10.1

Der Widerstand a hat den Nennwert

$$56 \cdot 10^3 \Omega, \text{ also } 56 \text{ k}\Omega$$

bei einer Toleranz von $\pm 10 \%$.

Für den Widerstand b ergibt sich ein Nennwert von $10 \cdot 10^4 \Omega = 100 \text{ k}\Omega$ bei 20 % Toleranz.

Beispiel 2:

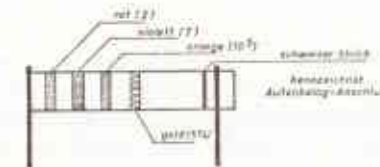


Abb. 10.2

Der Kondensator hat eine Kapazität von $27 \cdot 10^3 \text{ pF} = 27\,000 \text{ pF} = 27 \text{ nF}$ bei einer Toleranz von $\pm 5 \%$.

Für die Auswahl bestimmter Widerstandswerte wird empfohlen:

- den Nennwert bestimmen, in dessen Toleranzbereich der gesuchte Widerstand liegen muß und
- aus mehreren Widerständen des gleichen Nennwerts durch Messung mit dem Ohmmeter den Widerstand auswählen, dessen Wert dem gewünschten Wert am nächsten liegt.

10.4. Normreihen für Widerstände und Kapazitäten

Die Fertigung von Bauteilen unterliegt gewissen Streuungen. Es ist wirtschaftlich daher nicht vertretbar, z.B. Widerstände und Kondensatoren genau auf die vorgesehenen Werte abzugleichen. Ein solcher Abgleich wird nur für besondere Bauteile — z.B. für Meßwiderstände — vorgenommen und führt zwangsläufig zu einem erheblich höheren Preis. Man hat daher sog. Normreihen aufgestellt, die nach bestimmten Toleranzen gestuft sind. Die Toleranzgruppen sind $\pm 20\%$, $\pm 10\%$ und $\pm 5\%$ sowie für Sonderzwecke (Meßwiderstände) $\pm 1\%$ und $\pm 2\%$.

Würde man nun z.B. innerhalb der Toleranzgruppe 20% Widerstände mit von Ohm zu Ohm steigendem Wert herstellen, so ergäbe sich folgendes:

Nennwert	mögliche Abweichung	tatsächlich mögliche Werte
1 Ω	0,2 Ω	0,8 1,2 Ω
2 Ω	0,4 Ω	1,6 2,4 Ω
3 Ω	0,6 Ω	2,4 3,6 Ω
4 Ω	0,8 Ω	3,2 4,8 Ω
5 Ω	1,0 Ω	4,0 6,0 Ω
6 Ω	1,2 Ω	4,8 7,2 Ω
7 Ω	1,4 Ω	5,6 8,4 Ω

Wir erkennen, daß ein gesuchter Wert von 4,2 Ohm sowohl in der Nennwertgruppe 4 Ω als auch in der 5- Ω -Gruppe enthalten sein kann. Bei größeren Werten erstreckt sich die Streuung sogar über 3 oder mehr Nennwertbereiche (z.B. 4,8 ist möglich im Nennwertbereich 4 Ω , 5 Ω und 6 Ω).

Da sich also bei gleichmäßiger Erhöhung des Nennwerts die Bereiche infolge der Toleranz zunehmend überlappen, hat man aus wirtschaftlichen Gründen für Widerstands- und Kapazitätswerte sog. Normreihen aufgestellt. Die Normwerte sind dabei so festgelegt worden, daß sich innerhalb einer bestimmten Toleranz keine Überlappungen mehr ergeben. Die Reihen sind in folgender Tabelle zusammengestellt:

Internationale Normreihen

Toleranz:	20 %	10 %	5 %
Nennwerte	1,0	1,0	1,0
			1,1
		1,2	1,2
			1,3
	1,5	1,5	1,5
			1,6
		1,8	1,8
			2,0
	2,2	2,2	2,2
			2,4
		2,7	2,7
			3,0
	3,3	3,3	3,3
			3,6
		3,9	3,9
			4,3
	4,7	4,7	4,7
			5,1
		5,6	5,6
			6,2
	6,8	6,8	6,8
			7,5
		8,2	8,2
			9,1

Die in der Tabelle enthaltenen Nennwerte werden für hochohmige bzw. niederohmige Werte mit Zehnerpotenzen erweitert (z.B. Normwert 68 k Ω). An der Aufstellung erkennt man, daß der Nennwert bei großer Toleranz mehrere Nennwertbereiche für engere Toleranzen

umfaßt. Man kann also z.B. einen Widerstand von 56 k Ω in der 20-%-Toleranzgruppe unter dem Nennwert 47 k Ω finden. Wegen der großen Streuung muß dieser Widerstand durch Messen mit dem Ohmmeter ausgesucht werden.

10.5. Anschlußbeispiele für Dioden und Gleichrichter

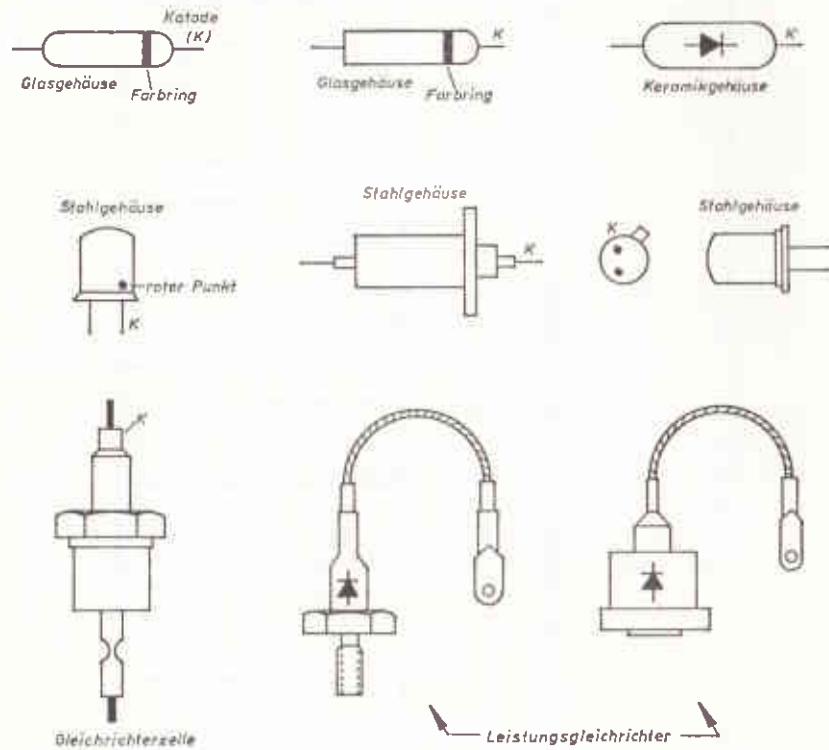


Abb. 10.3

10.6. Anschlußbeispiele für Transistoren

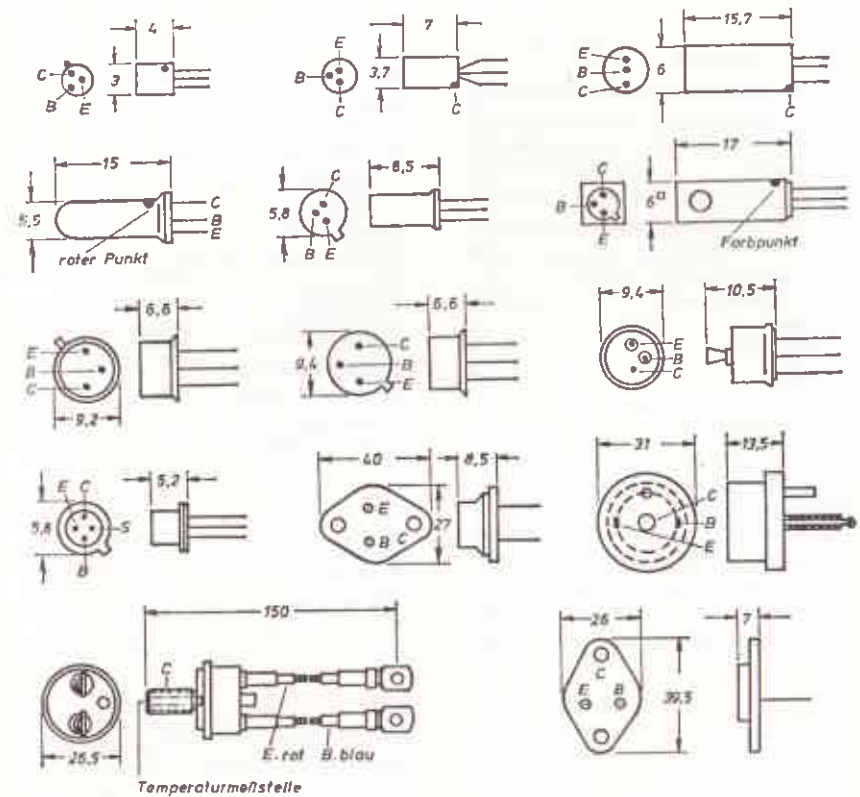


Abb. 10.4

Sachwortverzeichnis

A

Ablenkkoeffizient	27
akustische Rückkopplung	210
Anlaßheißleiter	56
Anlaufstrom (Röhre)	234, 239
Anpassung	153
Anschlußbeispiele	
— Diode und Gleichrichter	296
— Transistor	297
Anzeigegenauigkeit (Zeigerinstrument)	12
Arbeitspunkt	73
— Transistor	127
— Röhre	238
Arbeitspunktstabilisierung (Transistor)	167
astabile Kippstufe	254, 276
Ausgangskennlinien (Transistor)	141
Ausgangswiderstand (Transistor)	150, 158
Auskopplung des Signals	162
AUS-Stellung eines Transistor-schalters	203
Avalanche-Effekt	96

B

Balanceregler	196
Basis (Transistor)	123
Basis-Emitter-Spannung	140
Basisruhestrom	127, 149
Basisschaltung (Transistor)	156
Basis-Spannungsteiler (Berechnung —)	161
Basisvorwiderstand (Berechnung —)	181
Basisvorwiderstand (Berechnung —)	161
Berechnung —	180
Batteriebetrieb (Transistor)	159
Bedienplatte (Oszilloskop)	27
Bedienungselemente (Oszilloskop)	30
Beleuchtungsstärke (Fotodiode)	113
Berechnung	
— Basis-Spannungsteiler	181
— Basisvorwiderstand	180
— Kopplungskapazität	184
— Transistorverstärker	179

Beschleuniger

(RC-Glied als —)	248
Betriebsspannungen (Transistor)	160
Bezeichnung	
— Elektronenröhren	290
— Halbleiter-Bauelemente	288
— Spannungen und Ströme eines Transistors	131
bistabile Kippstufe	254, 256
bistabiler Multivibrator	254, 256
Blinkgeber (Selbstbau —)	277
Braunsche Röhre	21
Brückengleichrichter	82
Brückenschaltung	82
Brummsieb (Referenzdiode als —)	108

C

CDS-Fotowiderstand	64
CDS-Zelle	64
C-Kopplung	162

D

DIAC	228
differentieller Widerstand	73
Differenzglied als	243
— Hochpaß	250
— Impulsformer	244
Diffusion	68
Diffusionsspannung	68, 78
Diode	66
(Anschlußbeispiele —)	296
(Elektronenröhre als —)	234
Diode als	
— Entkoppler	87
— Gegenzelle	95
— Gleichrichter	78
— Leistungsbegrenzer	93
— Schalter	84
— Spannungsbegrenzer	89
direkte Heizung (Röhre)	232
Doppelbasistransistor	218

drain (FET)	217
Dreipolröhre	234
Drosselspule	84
Dunkelwiderstand (Fotowiderstand)	64
Durchlaßbereich (Diode)	71
Durchlaßrichtung (PN-Übergang)	67
dynamischer Eingang (Kippstufe)	265
dynamischer Widerstand	71
— Referenzdiode	97

E

e-Funktion	242
Eigenleitfähigkeit	44
Eigenverbrauch (Zeigerinstrument)	11
Eingangskennlinie (Transistor)	138
Eingangswiderstand (Transistor)	150, 158
Einkopplung des Signals	162
EIN-Stellung eines Transistor-schalters	201
Einweggleichrichter	78
Einwegschaltung	79
Elektronenoptik	22
Elektronenpaarbindung	42
Elektronenröhre (Bezeichnung —)	232
Elektronenstrahlröhre	290
Emittor	21
— Transistor	123
— Unijunktion-Transistor	219
Empfindlichkeit	
— Fotodiode (spektral)	113
— Zeigerinstrument	11
Entkoppler (Diode als —)	87

F

Farbcode für Widerstände und Kondensatoren	292
Feldeffekt-Transistor (FET)	216
Feldplatte	60
Flimmern (Oszilloskop)	32
Flipflop	256
Flipflop als	
— Frequenzteiler	269
— Speicher	269
— Zähler	271
FOCUS (Oszilloskop)	22

Folgefrequenz (Oszilloskop)	25
Formelzeichen für Spannung und Strom	11
Fotodiode	112
Fototransistor	220
Fotowiderstand	63
Frequenzbereich (Zeigerinstrument)	21
Frequenzmessung (Oszilloskop)	40
Frequenzteiler (Flipflop als —)	269
Frequenzverdopplung	82
Frequenzverhalten (Transistor)	154
Funkenlöschung (Varistor)	111
Funktionsprüfung	
— Diode	76
— Thyristor	225
— Transistor	132, 205

G

Galliumarsenid	115
Galliumarsenphosphid	115
Galliumphosphid	115
galvanische Kopplung	165
gate (FET)	217
Gegenkopplung	211
Gegenzelle (Diode als —)	95
Gehörschutzgleichrichter	91, 228
Getter (Röhre)	234
Gitter (Röhre)	235
Gitterspannungsquelle	238
Gittervorspannung	238
Gleichrichter (Diode als —)	78
(Röhre als —)	234
Gleichspannungsmessung	
— Oszilloskop	35
— Zeigerinstrument	14
Gleichstrom	9
Gleichstromkopplung (Transistor)	165
Gleichstrommessung	
— Oszilloskop	35
— Zeigerinstrument	15
Gleichstromverstärkungsfaktor (Transistor)	127, 136
Grenzfrequenz	
— RC-Glied	250
— Transistor	154, 158
Grenztemperatur (Diode)	78
Güteklasse (Zeigerinstrument)	12

H

Halbleiter	42
Halbleiter-Bauelement (Bezeichnung —)	288
Halbleiterdiode	66
Halbleiterstoff	42
Hallgenerator	62
Hallplatte (Hallsonde)	60
Haltespannung (Vierschichtdiode)	223
Haltestrom (Vierschichtdiode)	223
Heißeiter	53, 172
Helligkeit (Ozilloskop)	22
Helligkeitssteuerung (Lumineszenz- diode)	116
Hilfsspannung (Kippstufe)	260
Hinlaufzeit (Ozilloskop)	24
Hochpaß (RC-Glied als —)	249
Horizontalablenkung	23
HORIZONTAL INPUT (Ozilloskop)	23
HORIZONTAL AMPLITUDE (Ozilloskop)	24

I

Indirekte Heizung (Röhre)	232
indirekte Strommessung — Ozilloskop	35, 39
— Zeigerinstrument	19
Innenwiderstand — Ohmmeter	18
— Ozilloskop	31
— Spannungsmesser	14
— Strommesser	16
Integrierglied (RC-Glied als —)	245
Integrierglied — Kippschaltung	247
— Tiefpaß	252
INTENSITY (Ozilloskop)	22
internationale Normreihen	295
Impuls	10
Impulsformung — Differenzierglied	244
— Dioden	92
— Integrierglied	246
— Schmitt-Trigger	272
Impulsfrequenz (Kippstufe)	276, 278
Impulsgenerator (UJT)	220

J

Japan-Bezeichnungsschlüssel für Halbleiter-Bauelemente	289
---	-----

K

Kanalstrom (FET)	217
Kapazität (Auf- und Entladevorgang)	241
(Normreihen —)	294
Kapazitätsdiode	119
Katode — Diode	74
— Röhre	234
Katodenwiderstand (Röhre)	239
Kennlinie — DIAC	228
— Diode	70, 76
— NTC-Widerstand	55
— PTC-Widerstand	59
— Referenzdiode	97, 100
— Röhre	234, 236
— Transistor	138
— TRIAC	230
— Varistor	110
— Vierschichtdiode	223
— Widerstand	51
Kennliniendarstellung (Ozilloskop)	39
Kennlinienfeld	52
Kippschaltung	247
Kippstufe (Transistor)	254
Kippstufe (Selbstbau —)	277
Kleingleichrichter (Bezeichnungsschlüssel —)	290
Kollektor (Transistor)	123
Kollektor-Emitter-Rest- spannung	143, 203
Kollektorreststrom	206
Kollektorruhestrom	127, 149
Kollektor-Sättigungsspannung	143
Kollektorschaltung	156
Kompensations-Heißeiter	56
Kondensator (Farbcode —)	292
(Normreihen —)	294
(Siebketten —)	83

Kopplungsfaktor	210
Kopplungskapazität	184
Kristallgitter	42

L

Lastwiderstand (Bemessung —)	174, 183
LC-Generator	212
LED (lichtemittierende Diode)	115
Leistungsbegrenzer (Diode als —)	93
Leistungshyperbel	52
Leistungsverstärkerstufe	192
Leistungsverstärkung (Transistor)	153, 158
Lichtschranke	65, 119
Lichtsteuerung (Fotowiderstand)	63
Loch (Störstelle)	45
Löschen (Vierschichtdiode)	223
Lumineszenzdiode	115

M

magnetisch steuerbarer Wider- stand	61
Messungen — Diode	75
— Transistor	132
Mischstrom	10
Mitkopplung	211
Mittelwert, zeitlicher — Einweggleichrichtung	79
— Brückengleichrichtung	82
monostabile Kippstufe	254, 272
MOS-FET	218

N

negative Logik	256, 268
Netzbetrieb (Transistor)	159
Netzgerät, stabilisiert	200
NICHT-Glied	286
NIVEAU (Ozilloskop)	26
N-Leiter	44
Normreihen für Widerstände und Kapazitäten	294
NPN-Transistor	122

NTC-Widerstand	53, 172
Nullkippspannung (Vierschicht- halbleiter)	224
Nullpunktunterdrückung	107
Nullspannungsschalter	227

O

ODER-Glied	284
Ohmmeter	17
Oszillogramm	23
Oszillograf (Ozilloskop)	21

P

Parallelstabilisierung einer Spannung	197
Periodendauer (Messung —)	41
Phasenanschnittsteuerung	226
P-Leiter	44
PNP-Transistor	122
PN-Übergang	66
— Transistor	123
Polarität — Spannungsmessung	15
— Strommessung	16
— Widerstandsmessung	17
positive Logik	256
PTC-Widerstand	57
pulsierender Gleichstrom	10, 80
Pulsstrom	10

R

RC-Generator	213
RC-Glied als — Beschleuniger	248
— Differenzierglied	243
— Hochpaß	250
— Impulsformer	241
— Integrierglied	245
— Tiefpaß	252
RC-Kopplung	162
RC-Phasenkettengenerator	215
Rechteckgenerator (elektronischer Schalter als —)	32
Referenzdiode	96
Regelteil eines Transistor- verstärkers	193

Rekombinieren	68
Rückkopplung	209
Rücklaufzeit	24

S

Sägezahnimpuls	24, 26
Sägezahnspannung	9, 24
Sättigungsbereich (Röhre)	234
Schärfe (Oszilloskop)	22
Schalter	84
(Diode als —)	
(Transistor als —)	201, 255
Schmitt-Trigger	272
Schwellwertschalter	275
Schwingungserzeuger	
(Transistor als —)	209
Selbstbau einer Kippstufe	277
Serienstabilisierung einer Spannung	198
Siebensegment-Ziffernanzeige	117
Siebkette	83
Signalumkehrung	255
Signalzustände beim Transistor-schalter	255
Sonderformen des Transistors	216
source (FET)	217
Spannungen eines Transistors	131
Spannungsgegenkopplung	171, 185
Spannungsmessung	
— Oszilloskop	31
— Zeigerinstrument	14
Spannungsbegrenzer	
(Diode als —)	89
(Referenzdiode als —)	107
Spannungsstabilisierung	
— Diode	91, 173
— Referenzdiode	99
— Transistor	197
Spannungsverstärkung	
(Transistor)	152, 158
Speicher	
(Flipflop als —)	269
spektrale Empfindlichkeit	
(Fotodiode)	113
Spektralkurven (Lumineszenz-diode)	116
Sperrbereich (Diode)	71
Sperrichtung (PN-Übergang)	67
Sperrspannung (Diode)	68, 78

Spitzenspannungsmesser	247
stabilisiertes Netzgerät	200
stationärer Zustand	
— NTC-Widerstand	54
— PTC-Widerstand	59
statischer Eingang (Kippstufe)	261
statischer Widerstand	71
Steilheit (Röhre)	236
steuerbarer Gleichrichter	226
Steuerelektrode (Thyristor)	224
Steurgitter (Röhre)	235
Steuerkennlinien (Transistor)	144
Steuerstrom (Thyristor)	224
Störstellenleiter	44
Störstellenleitung	46
Strombelastbarkeit (Diode)	78
Stromgegenkopplung	169, 186
Strommessung	
(indirekte —)	19
Strommessung	
— Oszilloskop	31
— Zeigerinstrument	15
Stromsteuerung (Transistor)	146
Stromversorgung (Transistor-verstärker)	159
Stromversorgungsteil eines Transistorverstärkers	195
Stromverstärkung (Transistor)	151, 158
Ströme eines Transistors	131
Synchronisation (Oszilloskop)	25

T

Temperaturfühler	60
Temperaturverhalten (Halbleiter)	46
Thermistor	53
Tiefpaß	
(RC-Glied als —)	252
TIME BASE (Oszilloskop)	26
Thyristor	223
Transformatorkopplung	164
Transistor	122
(Anschlußbeispiele —)	297
(Sonderformen —)	216
Transistor als	
— Schalter	201, 255
— Schwingungserzeuger	209

— Verstärker	148
Transistorverstärker	
(Berechnung —)	179
(mehrstufige —)	192
Treiberstufe	192
TRIAC	229
Triggerdiode	228
Triggern (Oszilloskop)	25
Triggeroszilloskop	28
Triode (Röhre)	234

U

Überlastungsschutz (DIAC)	229
Überspannungsschutz	
(Varistor)	112
Übertragerkopplung	164
UND-Glied	282
Unijunktion-Transistor	218
Universaldiode	78
USA-Bezeichnungsschlüssel	289

V

Valenzbrücke	43
Valenzelektron	42
Varaktor	119
Varistor	109
VDR-Widerstand	109
Verknüpfungsglied	282
Verlustleistung	48, 50
Verstärker	
(Röhre als —)	237
(Transistor als —)	148
Verstärkerröhre	236
Verstärkungsfaktor	
— Röhre	238
— Transistor	126, 127, 136
Vertikalablenkung	23
VERTICAL AMPLITUDE	
(Oszilloskop)	26
VERTICAL INPUT (Oszilloskop)	23
Vielfachinstrument	12
Vierschichtdiode	222
Vierschichthalbleiter-Bauelemente	222
Vorverstärker	
(Aufbau —)	189
Vorverstärkerstufe	192

W

Wärmeableitung	50
Wehneltzylinder	22
Wechselspannungsmessung	
— Oszilloskop	36
— Zeigerinstrument	20
Wechselstrom	9
Wechselstrommessung	
— Oszilloskop	39
— Zeigerinstrument	20
Wechselstromverstärkungsfaktor	
(Transistor)	126
Widerstandsmessung mit dem	
Ohmmeter	17
Widerstände	
(Farbcode —)	292
(Normreihen —)	294
Wien-generator	214
Wien-Robinson-Brücke	213
Wirkungsgrad (Diode)	78

X

X-Ablenkung	23
X-Eingang	23
X-Position	24

Y

Y-Ablenkung	23
Y-Eingang	23
Y-Position	27

Z

Zähler	
(Flipflop als —)	271
Zählrichtung für Spannungen und Ströme (Transistor)	131
Z-Diode	96
Zeitablenkung (Oszilloskop)	24
Zeitbasis (Oszilloskop)	26
Zeitkonstante (RC-Glied)	243
Zenerbereich	96
Zenerdiode	96
Zenereffekt	96
Zenerspannung	96
Zweipolröhre	234
Zweiweggleichrichter	81
Zweiwegschaltung	81
Zündspannung (Vierschichtdiode)	222

Teil 3 — Datenverarbeitung; Technik und Betrieb

Nachdem im ersten und zweiten Teil des „Handbuchs der Elektronik“ die Grundlagen der Elektronik und ihre Anwendung in der linearen und der digitalen Technik eingehend behandelt worden sind, werden diese Themen im dritten Teil weitergeführt. Im einzelnen werden folgende Stoffgebiete ausführlich behandelt:

Arbeitsweise und Organisation von EDV-Anlagen
Leitwerk
Rechenwerk
Speicher
Geräte der Peripherie
Einführung in die Mengenlehre
Grundzüge der Aussagenlogik
Wesen der Schaltalgebra
Boolesche Algebra

Das Lehr- und Lernbuch enthält eine Vielzahl von Abbildungen und viele ausführliche Rechenbeispiele, die für den Techniker auf praktische Fälle zugeschnitten sind.

Fachwörter der Elektronik

— Zur Ergänzung des Lehrprogramms Halbleitertechnik und Elektronik —

In diesem Nachschlagewerk, das in enger Beziehung zu den vorgenannten Lehr- und Lernbüchern steht, sind alle wichtigen Fachwörter und Begriffe der Halbleitertechnik und der Elektronik in alphabetischer Reihenfolge übersichtlich geordnet zusammengestellt. Zum besseren Verständnis werden die Stichwörter durch Erläuterungen und Erklärungen sinnvoll ergänzt. Hier können die vielen, oft für den Leser neuen und noch nicht fest in sein Bewußtsein eingeordneten technischen Begriffe oder Bezeichnungen schnell nachgelesen und Unklarheiten rasch beseitigt werden.

Unabhängig von dieser Zielsetzung findet der bereits versierte Elektroniker in dem Fachwörterverzeichnis viele Erklärungen, die ihm das Lesen der Funktionsbeschreibungen für die im praktischen Betrieb eingeführten elektronischen Anlagen ganz wesentlich erleichtern werden.

Tabellenbuch der Nachrichtentechnik

Das neue Nachschlagewerk umfaßt rd. 290 Seiten mit mehr als 800 Abbildungen und erscheint im Dreifarbendruck (Format DIN A5); es gliedert sich in folgende große Abschnitte:

1. Fernmelderechnen,
2. Fernmeldebautelle und
3. Fernmeldesachbereiche.

Ein Überblick über sämtliche in der Fernmeldetechnik verwendeten Schaltzeichen schließt das Tabellenbuch ab. Das schnelle Auffinden der nachzuschlagenden Tabellen, Wertangaben, Daten usw. ist anhand eines ausführlichen Sachwortverzeichnisses leicht möglich. Bei der Stoffdarstellung ist weitgehend auf DIN-Vorschriften sowie die VDE- und PTZ Vorschriften zurückgegriffen worden.

Gesamtübersicht über das Lehrmittelvorhaben Elektrotechnik Halbleitertechnik Elektronik Datenverarbeitung

- **Grundlagen der Elektronik**
 - Repetitor „Grundlagen der Elektronik“

Handbuch der Elektronik

- **Teil 1 — Analogtechnik**
 - Repetitor „Analogtechnik“
 - **Teil 2 — Digitaltechnik**
 - Repetitor „Digitaltechnik“
 - **Teil 3 — Datenverarbeitung; Technik und Betrieb**
 - **Fachwörter der Elektronik**
 - **Tabellenbuch der Nachrichtentechnik**
-

Sämtliche Lehrwerke können bestellt werden bei dem

Institut zur Entwicklung moderner Unterrichtsmethoden e. V.

28 Bremen 1, Bahnhofstraße 10

Fernsprecher 0421 / 31 52 85