

„Elektronische Schaltungen“

Karlsruhe, Sommersemester 2018

Prof. M. Siegel, Dr. S. Wunsch

Postanschrift:	Institut für Mikro– und Nanoelektronische Systeme	Tel.:	+49 (0) 721 608 4 49 61
	Hertzstraße 16	Sekr.:	+49 (0) 721 608 4 49 61
	D - 76128 Karlsruhe	Fax.:	+49 (0) 721 75 79 25
		Email:	doris.duffner@kit.edu
Gebäude:	Hertzstraße 16, Geb. 06.41	WWW:	http://www.ims.kit.edu

Inhaltsverzeichnis

1	Allgemeine Grundlagen	2
1.1	Definitionen und Formelzeichen	2
1.2	Grundelemente und Schaltungssymbole	3
1.3	Die Kirchhoff'schen Gesetze	4
1.3.1	Knotenregel	4
1.3.2	Maschenregel	4
2	Passive Bauelemente	5
2.1	Allgemeines über Bauelemente	5
2.2	Kennwerte und Kennzeichnungen von Bauelementen	6
2.3	Ohm'sche Widerstände	10
2.3.1	Bauformen von Widerständen	12
2.3.2	Temperaturabhängigkeit von Widerständen	12
2.3.3	Belastbarkeit von Widerständen	13
2.3.4	Nichtlineare Widerstände	14
2.4	Kondensatoren	16
2.4.1	Bauformen von Kondensatoren	17
2.5	Spulen	19
2.5.1	Bauformen	20
3	Halbleiterbauelemente	21
3.1	Die Diode	22
3.1.1	Der Aufbau	22
3.1.2	Die Kennlinien	23
3.1.3	Ersatzschaltbild und statisches Verhalten	26
3.1.4	Das Kleinsignalverhalten	27
3.1.5	Das dynamische Verhalten	28
3.1.6	Vollständiges Modell einer Diode	29
3.1.7	Vereinfachtes dynamisches Kleinsignal-Modell	31
3.1.8	Anwendungsbeispiele	31
3.1.9	Die Zener-Diode	33
3.2	Bipolare Transistoren	35
3.2.1	Aufbau und Ersatzschaltbild	35
3.2.2	Die Funktionen eines Bipolartransistors	36
3.2.2.1	Statische Kennlinien	36
3.2.2.2	Der Early-Effekt	40
3.2.3	Arbeitspunkt und Kleinsignal-Verhalten	41
3.2.3.1	Bestimmung des Arbeitspunktes	41

3.2.3.2	Die Kleinsignal-Parameter	42
3.2.3.3	Wechselstrom–Kleinsignal–Ersatzschaltbild	45
3.2.3.4	Grenzdaten und zuverlässiger Arbeitsbereich	46
3.2.4	Die Grundsaltungen	48
3.2.4.1	Die Emitterschaltung	49
3.2.4.2	Die Kollektorschaltung	59
3.2.4.3	Die Basisschaltung	63
3.2.4.4	Kleinsignalverhalten der Basisschaltung	64
3.3	Sperrschicht–Feldeffekttransistoren (JFET)	67
3.3.1	Aufbau und Wirkungsweise	67
3.3.2	Funktionen eines Sperrschicht–Feldeffekttransistors	69
3.3.2.1	Eingangs– und Ausgangskennlinien	69
3.3.2.2	Der Early–Effekt	71
3.3.2.3	Bestimmung der Schwellspannung U_{th}	72
3.3.3	Arbeitspunkt und Kleinsignalverhalten	73
3.3.3.1	Kleinsignal–Ersatzschaltbild	75
3.3.3.2	Grundsaltungen mit Sperrschicht–Feldeffekttransistoren	76
3.4	Isolierschicht–Feldeffekttransistoren (MOSFET)	80
3.4.1	Aufbau und Wirkungsweise	80
3.4.2	Funktion und Kennlinien eines MOSFET	82
3.4.3	Die Grundsaltungen	86
3.4.3.1	Die Sourceschaltung	86
3.4.3.2	Die Drainschaltung	93
3.4.3.3	Die Gateschaltung	95
3.5	Schaltungen mit komplementären Feldeffekttransistoren (CMOS)	98
3.5.1	Aufbau und Wirkungsweise	98
3.5.2	Funktion und Kennlinien	99
3.5.3	Die Grundsaltungen	103
4	Verstärkerschaltungen	105
4.1	Mehrstufige Verstärker	106
4.2	Stromquellen und Stromspiegel	108
4.2.1	Stromquellen	108
4.2.1.1	Stromquellen mit diskreten Transistoren	110
4.2.2	Der Stromspiegel	111
4.3	Der Differenzverstärker	112
4.4	Gegentaktverstärker	116
4.5	Rückgekoppelte Verstärker	119
4.5.1	Gegenkopplung	120
4.5.2	Mitkopplung	121
5	Operationsverstärker	122
5.1	Aufbau und Wirkungsweise	122

5.2	Grundsaltungen mit Operationsverstärkern	126
5.2.1	Der invertierende Verstärker	126
5.2.2	Der nichtinvertierende Verstärker	127
5.3	Offsetspannung und Offsetstrom	128
5.4	Eingangs- und Ausgangswiderstand	130
5.5	Frequenzgang	132
5.6	Gegengekoppelte Schaltungen mit Operationsverstärkern	136
5.6.1	Invertierender Integrator	136
5.6.2	Invertierender Differenzierer	137
5.6.3	Addierer	138
5.6.4	Subtrahierer	138
5.6.5	Logarithmierer	140
5.6.6	Exponentialfunktion	142
5.6.7	Messverstärker	143
6	Kippschaltungen	145
6.1	Kippschaltungen mit bipolaren Transistoren	145
6.1.1	Bistabile Kippschaltungen	146
6.1.1.1	Flip-Flop	146
6.1.1.2	Schmitt-Trigger	147
6.1.1.3	Monostabile Kippschaltungen	151
6.1.1.4	Astabile Kippschaltungen	152
6.1.2	Kippschaltungen mit Komparatoren	153
6.2	Kippschaltungen mit Komparatoren	153
6.2.1	Invertierender Schmitt-Trigger	153
6.2.2	Nichtinvertierender Schmitt-Trigger	154
6.3	Multivibratoren mit Operationsverstärker	155
7	Grundlagen digitaler Schaltungen	158
7.1	Inverter und Störabstände	158
7.2	Anstiegs- und Gatterlaufzeiten	160
7.3	Die Verlustleistung	161
7.4	Lastfaktoren	162
7.5	Positive und negative Logik	163
7.6	Genormte Schaltzeichen	163
7.6.1	Darstellung von Eingängen	163
7.6.2	Darstellung von Ausgängen	164
7.6.3	Logische Grundelemente	166
7.6.4	Beispiele	170
8	Schaltkreisfamilien	172
8.1	Bipolare integrierte Schaltungen	172
8.1.1	DTL-Schaltungen	172
8.1.2	TTL-Schaltungen	173

8.1.3	Schottky- und Low-Power Schottky-TTL-Schaltungen	177
8.1.4	Übersicht über die Schaltkreisfamilien	182
8.2	MOS-Schaltkreise	184
8.2.1	NMOS-Schaltungen	184
8.2.2	CMOS-Schaltungen	187
8.2.3	Anpassungsschaltungen	191
8.2.4	CMOS-Baureihe 4000	192
8.2.5	Hochgeschwindigkeits-HC/HCT-Reihe	193
8.2.6	Advanced CMOS-Reihe AC/ACT	194
8.2.7	Advanced High-Speed-CMOS-Reihe AHC/AHCT	195
8.2.8	BiCMOS-Schaltungen	196
8.2.9	Elektronische Schalter (Transferrgatter, Transmission Gate)	200
9	Sequentielle Logik	202
9.1	Klassifizierung	202
9.1.1	Zustandsgesteuerte Flip-Flop	203
9.1.2	Taktzustandsgesteuerte Flip-Flop	207
9.1.3	Taktflankengesteuerte Flip-Flop	215
9.2	Zähler	223
9.2.1	Asynchrone Zähler	224
9.2.2	Asynchrone Dualzähler	224
9.2.2.1	Asynchrone Modulo-n-Zähler	231
9.2.2.2	Synchrone Zähler	234
9.2.3	Zähler als Frequenzteiler	241
9.2.3.1	Asynchrone Frequenzteiler mit festem Teilverhältnis	241
9.2.3.2	Synchrone Frequenzteiler mit festem Teilverhältnis	244
9.2.3.3	Synchrone Frequenzteiler mit einstellbarem Teilverhältnis	245
9.3	Schieberegister	246
9.3.1	Schieberegister für serielle Ein- und Ausgabe	246
9.3.2	Schieberegister mit Parallelausgabe	250
9.3.3	Schieberegister mit Parallelausgabe und Paralleleingabe	251
9.3.4	Ringregister	252
10	Decoder (Codewandler, Dekodierer), Multiplexer	254
10.1	n bit binär zu 1-aus-n Decoder	254
10.2	BCD zu 7-Segment Decoder	254
10.3	Prioritätsdecoder	254
10.4	Multiplexer, Demultiplexer	258
11	Digital-Analog- und Analog-Digital-Wandler	261
11.1	Digital-Analog-Wandler	261
11.1.1	Grundlagen der Digital-Analog Wandlung	261
11.1.2	Digital-Analog Wandlung nach dem Wägeverfahren	263

11.1.3 Genauigkeit von Digital–Analog–Wandlern	266
11.2 Analog–Digital Wandlung	268
11.2.1 Grundlagen der Analog–Digital Wandlung	268
11.2.2 Parallel–Wandler (Flash)	271
11.2.3 Analog–Digital Wandlung mit sukzessiver Approximation	273
11.2.4 Analog–Digital Wandlung mit Integrationsverfahren	275
11.2.5 Genauigkeit von A/D–Wandlern	276
12 Literatur	279
13 Formelsammlung zur Klausur “Elektronische Schaltungen”	280

Vorwort

Jeder Ingenieur der Elektrotechnik und Informationstechnik, gleich welcher Vertiefungsrichtung, sollte heute in der Lage sein, mehr oder weniger komplexe elektronische Schaltungen selbst zu entwickeln. Dazu ist es notwendig, die Eigenschaften der einzelnen Bauelemente und Grundsaltungen zu kennen. Eine elektronische Schaltung entsteht im Allgemeinen durch eine sinnvolle Verbindung mehrerer Bauelemente entweder auf einer Platine oder aber direkt in einer integrierten Schaltung.

Ziel dieser Vorlesung ist es, Kenntnisse über die Eigenschaften der wichtigsten Bauelemente zu vermitteln und mit diesen Bauelementen wiederum die gebräuchlichsten Grundsaltungen zu berechnen, um dann mit Hilfe dieser Kenntnisse funktionierende Schaltungen entwerfen zu können. Mit diesen Kenntnissen sollte es darüber hinaus auch möglich sein, komplexere Probleme anzugehen um zu einer, der Problemstellung angemessenen, Lösung zu kommen, die der Elektronik je nach Anforderung entweder nur mit analogen, nur mit digitalen oder einer Kombination von analogen und digitalen Schaltungen gelöst werden kann.

Bei den Hörern dieser Vorlesung wird die Kenntnis der Mathematik und der Physik der Oberstufe der Gymnasien, des Stoffes der höheren Mathematik der ersten beiden Semester und der Vorlesung lineare elektrische Netze vorausgesetzt. Wichtig ist vor allem die Fähigkeit mit Zahlen und Einheiten zu rechnen. Die Kenntnis der Physik der verwendeten Bauelemente ist von Vorteil, jedoch nicht unbedingt erforderlich, da die wichtigsten physikalischen Eigenschaften, vor allem die der Halbleiterbauelemente, in dieser Vorlesung behandelt werden.

Die Entwicklung moderner elektronischer Geräte, Baugruppen oder integrierter Schaltungen erfordert neben theoretischen Kenntnissen auch eine gewisse praktische Erfahrung, die im Rahmen dieser Vorlesung jedoch nicht vermittelt werden kann. Diese praktischen Erfahrungen müssen durch die Teilnahme an Praktika oder im Rahmen einer Bachelor- oder Masterarbeit gesammelt werden.

Jedes Kapitel beginnt mit der Nennung der Lernziele. Die wichtigsten Formeln sind Bestandteil der Formelsammlung, die bei der Prüfung verwendet werden kann.

1 Allgemeine Grundlagen

Lernziele:

- Verstehen von Definitionen, Formelzeichen und Einheiten elektrischer Größen
- Arbeiten mit Grundelementen und Symbolen in der Schaltungstechnik
- Wiederholung der Methoden zur Berechnung von linearen elektrischen Netzen

Zum allgemeinen Verständnis des gesamten Bereiches moderner elektronischer Schaltungen ist eine, für alle verständliche, gleiche Darstellung der Bezeichnungen und Begriffe von großem Vorteil. Deshalb sollen zuerst die wichtigsten Definitionen und Formelzeichen, eine genormte graphische Darstellung der für die Schaltungstechnik notwendigen Elemente und einige immer wiederkehrende Rechenregeln und Bezeichnungen kurz zusammengefasst werden. Im weiteren Verlauf der Vorlesung werden weitere Definitionen und Symbole hinzukommen, die eingesetzt werden, um die Schaltungen darzustellen.

1.1 Definitionen und Formelzeichen

In dieser Vorlesung wird eine Reihe von Bezeichnungen verwendet, die im speziellen Fall dann auch näher erläutert werden. Immer wiederkehrende Begriffe und Formelzeichen und die zugehörigen Einheiten sind hier zunächst zusammengefasst.

Gleichstrom, Gleichspannung	$I[A], U[V]$
Wechselstrom, Wechselspannung	i, u
Leistung	$P [W=V \cdot A]$
Widerstand	$R [\Omega=V/A]$
Differentieller Widerstand	$r = \partial u / \partial i [\Omega = V/A]$
Leitwert	$G [1/R = A/V]$
Kapazität	$C [F = As/V]$
Induktivität	$L [H = Vs/A]$
Frequenz	$f [Hz = 1/s]$
Kreisfrequenz	$\omega = 2 \cdot \pi \cdot f [Hz]$
Zeitkonstante	$\tau [s]$
Stromverstärkungsfaktor (bipolare Transistoren)	B, β
Temperaturspannung	$U_T [V]$
Schwellspannung	$U_{th} [V]$
Steilheit	$S, g_m [A/V]$

1.2 Grundelemente und Schaltungssymbole

Die wichtigsten Grundelemente der Elektrotechnik zum Aufbau und zur Analyse von Schaltungen sind: Widerstand, Kapazität, Induktivität, Strom- und Spannungsquelle.

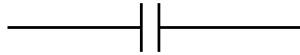
Widerstand R :



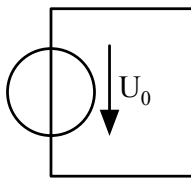
Induktivität L :



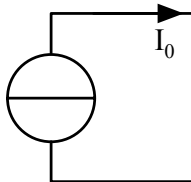
Kapazität C :



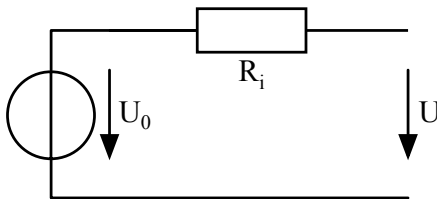
ideale Spannungsquelle:



ideale Stromquelle:



reale Spannungsquelle:



reale Stromquelle:

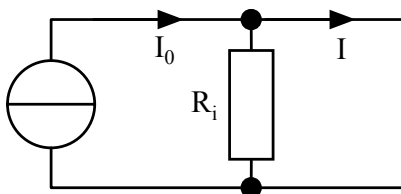


Bild 1.1: Übersicht der Grundelemente und Schaltungssymbole.

1.3 Die Kirchhoff'schen Gesetze

Zur Berechnung des Verhaltens der Schaltungen werden immer wieder die folgenden Grundregeln benötigt werden:

1.3.1 Knotenregel

Die Summe aller Ströme an einem Knoten ist Null (zufließende Ströme erhalten ein positives Vorzeichen, abfließende ein negatives).

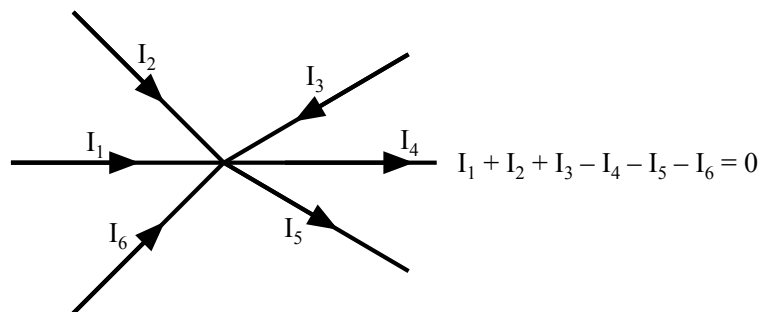


Bild 1.2: Darstellung der Knotenregel.

1.3.2 Maschenregel

Die Summe aller Teilspannungen eines geschlossenen Umlaufs einer Masche ist Null.

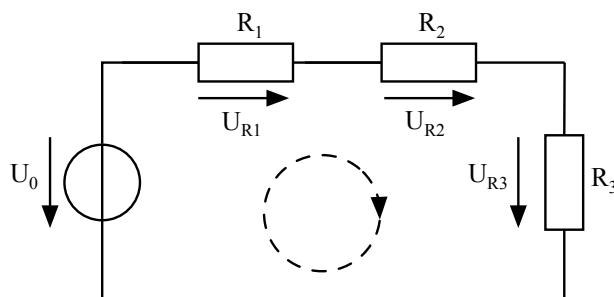


Bild 1.3: Darstellung der Maschenregel.

2 Passive Bauelemente

Lernziele

- Grundsätzliche Betrachtung von passiven Bauelementen der Schaltungstechnik
- Kennenlernen der gebräuchlichen Kennzeichnungen und Kennwerte für passive Bauelemente
- Eigenschaften elektronischer passiver Bauelemente

2.1 Allgemeines über Bauelemente

Als passive Bauelemente werden in der Schaltungstechnik üblicherweise Widerstände, Kondensatoren und Induktivitäten bezeichnet. Die Kennwerte und weiteren Merkmale der Bauelemente werden normalerweise von den Herstellern in Datenblättern angegeben.

Im folgenden sollen einige wichtige Merkmale, die auch eine Aussage über die Qualität des Bauelements machen kurz erläutert werden.

Toleranz	Streuung des tatsächlichen Wertes des Bauelements gegenüber dem Nennwert (Angabe üblicherweise in %.)
Temperaturkoeffizient	Änderung des Kennwertes in Abhängigkeit von der Betriebs- und Umgebungstemperatur.
Linearität	Änderung des Kennwertes in Abhängigkeit von Strom und Spannung beim Betrieb des Bauelements.
Belastbarkeit	Maximalwerte von Strom, Spannung und Leistung.
Wärmewiderstand	Wärmeableitung im Bauelement (von innen nach außen).
Zuverlässigkeit	Einhaltung der Kennwerte über die mittlere Lebensdauer des Bauelements.

Parasitäre Eigenschaften	Induktivität der Anschlussdrähte, Kapazität zwischen den Anschlüssen bzw. innerhalb des Bauelements, Leckströme im Innern des Bauelements usw.
Rauschen	Überlagerung des Ausgangstromes bzw. der Ausgangsspannung mit einer Wechsellspannung mit einer sehr kleinen Amplitude aber einem sehr breiten Frequenzspektrum

2.2 Kennwerte und Kennzeichnungen von Bauelementen

Moderne Bauelemente werden heute mit fast beliebigen Kennwerten und einer sehr hohen Genauigkeit hergestellt, welche letztendlich durch die Genauigkeit der eingesetzten Messverfahren bei der Überprüfung und beim Abgleich des Wertes bestimmt wird. Eine extrem hohe Genauigkeit ist jedoch sehr kostenintensiv und wird nur dann eingesetzt, wenn es unbedingt erforderlich ist. Deshalb werden die meisten Bauelemente heute mit Kennwerten und Toleranzen hergestellt, die dem Verwendungszweck angepasst sind.

Die einzelnen Kennwerte sind so abgestuft, dass die obere Toleranzgrenze des ersten Wertes an die untere Toleranzgrenze des nächsten Wertes innerhalb einer Reihe von Kennwerten heranreicht. Die Abstufung innerhalb der Reihen wird nach den international üblichen E-Reihen vorgenommen. Demnach hat eine Reihe n Werte pro Dekade, wobei $n = 3, 6, 12, 24, 48$, oder 96 sein kann. Die Abstufung der Werte innerhalb einer E-Reihe erfolgt logarithmisch mit einem Stufenfaktor, der sich aus der n -ten Wurzel von 10 errechnet. Damit sind aber auch die Toleranzwerte einer jeden Reihe genau bestimmt. In Tabelle 2.1 sind die Werte der E-Reihen über eine Dekade mit Angabe der Toleranz aufgelistet.

Die Werte aller anderen Dekaden erhält man, in dem man die in der Tabelle angegebenen Werte mit positiven bzw. negativen Potenzen von 10 multipliziert. Die Kennzeichnung der Bauelemente mit ihrem Wert und ihrer Toleranz geschieht auf vielfältige Weise. Sind die Bauelemente groß genug, wird Wert und Toleranz direkt aufgedruckt. Durch die immer kleiner werdenden Abmessungen der Bauelemente ist dies nicht mehr möglich. Deshalb verwendet man verschiedene Arten der Codierung, von denen einige näher erläutert werden sollen. Eine auch heute noch übliche Codierung ist die Farbcodierung, entweder durch Punkte oder durch Ringe. Farbringe haben den Vorteil, dass Wert und Toleranz des Bauelements unabhängig von der Einbaulage auf der Platine abgelesen werden können. Die Kennzeichnung erfolgt nach DIN 41429 bzw. nach der internationalen IEC-Norm. Die einfachste Form dieser Farbcodierung ist der 4-Ring-Farbcode, wie in Bild 2.1 gezeigt.

E3	E6	E12	E24	E48	E96		E3	E6	E12	E24	E48	E96
±20%	±20%	±10%	±5%	±2%	±1%		±20%	±20%	±10%	±5%	±2%	±1%
1,0	1,0	1,0	1,0	1,00	1,00						3,16	3,16
					1,02							3,24
				1,05	1,05			3,3	3,3	3,3	3,32	3,32
					1,07							3,40
			1,1	1,10	1,10						3,48	3,48
					1,13							3,57
				1,15	1,15					3,6	3,65	3,65
					1,18							3,75
	1,2	1,2	1,2	1,21	1,21						3,83	3,83
					1,24				3,9	3,9		3,92
				1,27	1,27						4,02	4,02
			1,3		1,30							4,12
				1,33	1,33						4,22	4,22
					1,37					4,3		4,32
				1,40	1,40						4,42	4,42
					1,43							4,53
				1,47	1,47						4,64	4,64
	1,5	1,5	1,5		1,50		4,7	4,7	4,7	4,7		4,75
				1,54	1,54						4,87	4,87
					1,58							4,99
			1,6	1,62	1,62					5,1	5,11	5,11
					1,65							5,23
				1,69	1,69						5,36	5,36
					1,74							5,49
				1,78	1,78				5,6	5,6	5,62	5,62
		1,8	1,8		1,82							5,76
				1,87	1,87						5,90	5,90
					1,91							6,04
				1,96	1,96					6,2	6,19	6,19
			2,0		2,00							6,34
				2,05	2,05						6,49	6,49
					2,10							6,65
				2,15	2,15			6,8	6,8	6,8	6,81	6,81
2,2	2,2	2,2	2,2		2,21							6,98
				2,26	2,26						7,15	7,15
					2,32							7,32
				2,37	2,37					7,5	7,50	7,50
			2,4		2,43							7,68
				2,49	2,49						7,87	7,87
					2,55							8,06
				2,61	2,61				8,2	8,2	8,25	8,25
					2,67							8,45
		2,7	2,7	2,74	2,74						8,66	8,66
					2,80							8,87
				2,87	2,87					9,1	9,09	9,09
					2,94							9,31
			3,0	3,01	3,01						9,53	9,53
					3,09							9,76

Tabelle 2.1: E-Reihen über eine Dekade.

Für Werte aus den E48- und E96-Reihen werden zur Angabe jedoch 3 Stellen benötigt, eine Stelle mehr als beim in Bild 2.1 dargestellten 4-Ring-Code. Deshalb wird dort der 5-Ring-Farbcode nach IEC 62 verwendet, wie er in Bild 2.2 zu sehen ist. Weitere Codierungen von Bauelementen verwenden Ziffernfolgen ähnlich der 4-Ring bzw. der 5-Ring-Farbcodierung oder Kombinationen aus Buchstaben und Zahlen wie sie in Bild 2.2 und Tabelle 2.2 angegeben sind und vor allem bei den modernen SMD-Bauelementen verwendet werden.

A	1,0		M	3,0		Y	8,2		0		1,0
B	1,1		O	3,3		Z	9,1		1		10,0
C	1,2		P	3,6		a	2,5		2		100,0
D	1,3		Q	3,9		b	3,5		3		1.000,0
E	1,5		R	4,3		d	4,0		4		10.000,0
F	1,6		S	4,7		e	4,5		5		100.000,0
G	1,8		T	5,1		f	5,0		6		1.000.000,0
H	2,0		U	5,6		Y	6,0		7		10.000.000,0
I	2,2		V	6,2		Y	7,0		8		100.000.000,0
K	2,4		W	6,8		Y	8,0		9		0,1
L	2,7		X	7,5		Y	9,0				

Tabelle 2.2: Einfachster Code für SMD-Bauelemente. Buchstabe = Wert, Ziffer = Multiplikator. Angaben für Widerstände in Ω , für Kondensatoren in pF.

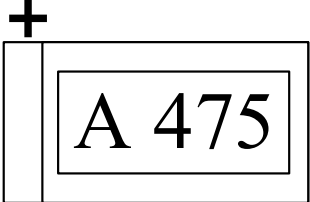
Buchstabe	Spannung [V]	Zahl = Kapazität
e	2,5	 Z.B. A475 ist 4,7 μF für 10 V 475 = $47 \cdot 10^5$ pF = $4,7 \cdot 10^6$ pF = 4,7 μF
G	4	
J	6,3	
A	10	
C	16	
D	20	
E	25	
V	35	
H	50	

Tabelle 2.3: Codierung für SMD-Elektrolytkondensatoren.

Für SMD Elektrolytkondensatoren wird üblicherweise die Codierung nach Tabelle 2.3 verwendet. Für SMD-Widerstände wird auch gerne ein 3- oder 4-stelliger Zifferncode verwendet, der mit dem Buchstaben R ergänzt wird, wenn die Widerstandswerte kleiner als 10 Ω sind. Sonst geben die ersten Stellen dabei immer den Widerstandswert an und die letzte Stelle den Multiplikator in Form des Exponenten zur Basis 10.

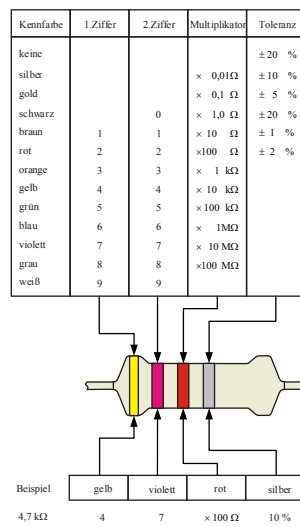


Bild 2.1: Übersicht der Grundelemente und Schaltungssymbole.

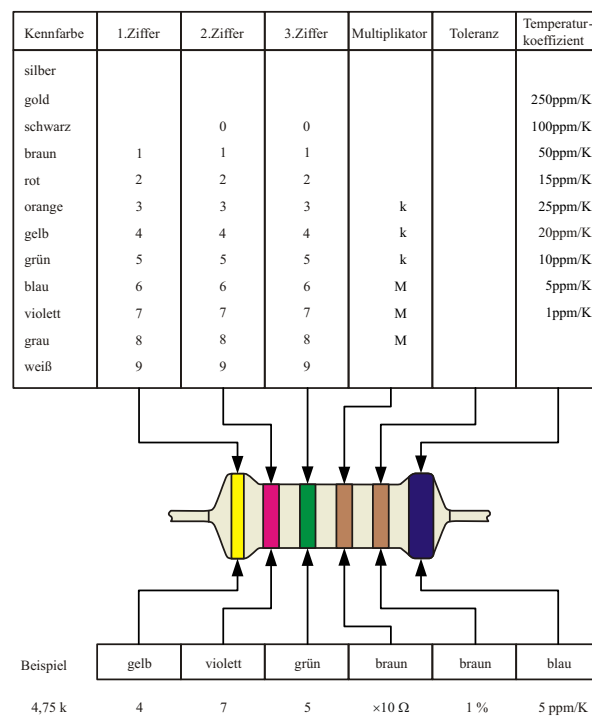


Bild 2.2: 5-Ring-Code nach IEC 62.

Beispiele:

3 Stellen:

$$330 = 33 \cdot 10^0 \Omega = 33 \Omega$$

$$684 = 68 \cdot 10^4 \Omega = 680 \text{ k}\Omega$$

$$8R2 = 8,2 \Omega$$

4 Stellen:

$$1000 = 100 \Omega$$

$$1623 = 162 \text{ k}\Omega$$

$$4992 = 49900 \Omega = 49,9 \Omega$$

$$0R56 \text{ oder } R56 = 0,56 \Omega$$

2.3 Ohm'sche Widerstände

Ein ohm'scher Widerstand ist ein Bauelement mit zwei Anschlüssen, das durch den Strom i durch das Bauelement in Abhängigkeit von der angelegten Spannung u beschrieben werden kann. Diese Abhängigkeit ist in der Form

$$u = R \cdot i \quad (2.1)$$

als Ohm'sches Gesetz bekannt. R hat die Dimension

$$[R] = \frac{[u]}{[i]} = \frac{V}{A} = \Omega \text{ (Ohm)} \quad (2.2)$$

und wird ohm'scher Widerstand genannt. Die lineare Abhängigkeit zwischen Strom und Spannung lässt sich in einem Strom–Spannungsdiagramm besonders einfach darstellen. Die I/U–Kennlinien von ohm'schen Widerständen sind Geraden mit der Steigung $\frac{1}{R}$ nach Bild 2.3. Der Wert eines Widerstands ist proportional dem spezifischen Widerstand ρ des verwendeten Materials, proportional der Länge ℓ des Drahtes oder der Schicht und umgekehrt proportional dem Querschnitt A , wie die auch die Beispiele in Bild 2.4 zeigen. Es ist also

$$R = \rho \cdot \frac{\ell}{A} \quad (2.3)$$

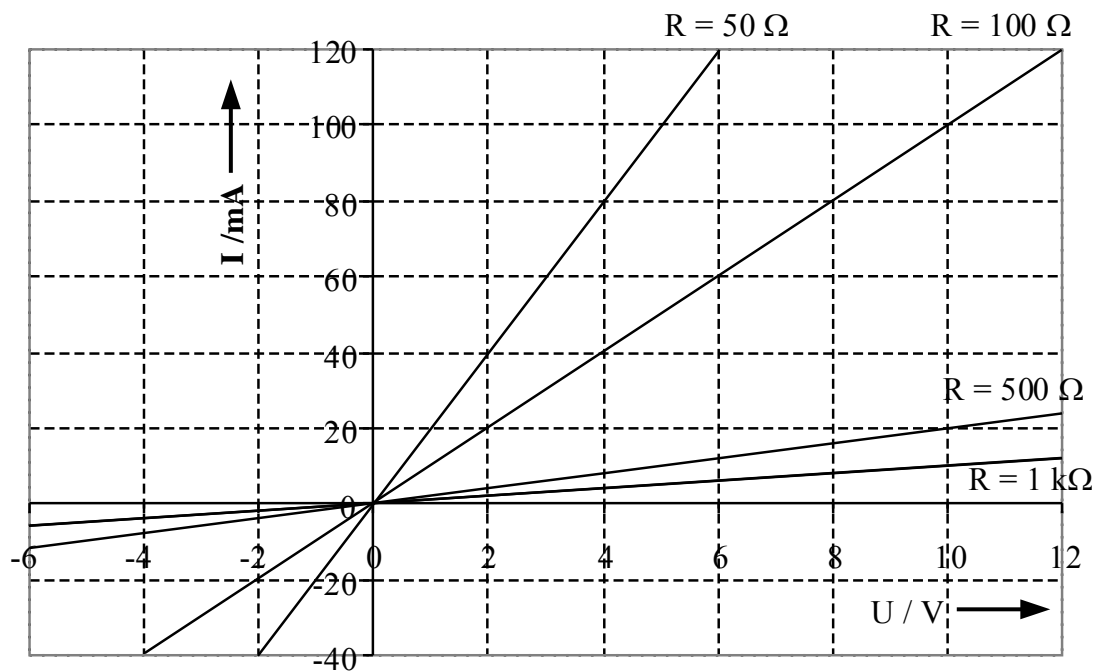


Bild 2.3: Kennlinien verschiedener ohm'scher Widerstände.

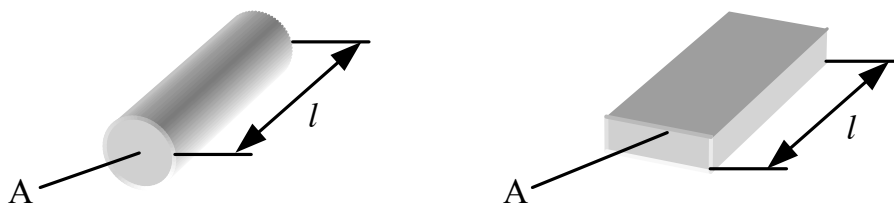


Bild 2.4: Schematische Darstellung von Leiterquerschnitten.

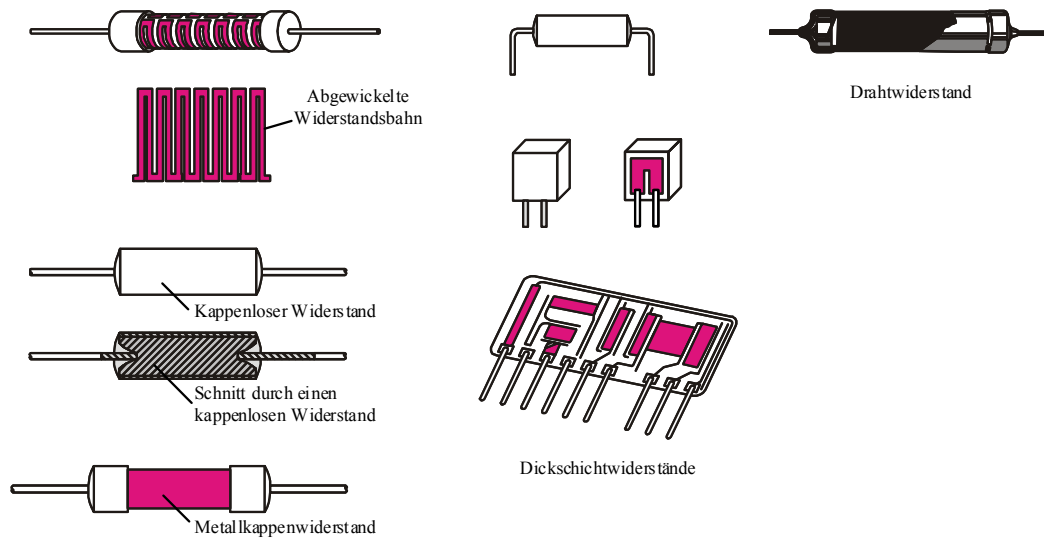


Bild 2.5: Bauformen von Widerständen.

2.3.1 Bauformen von Widerständen

Ohm'sche Widerstände werden in vielen verschiedenen Größen und Bauformen hergestellt, um den vielschichtigen Anforderungen Rechnung zu tragen und für alle Bereiche der Elektronik und Elektrotechnik eingesetzt zu werden. In Bild 2.5 sind einige Beispiele der verschiedenen Bauformen dargestellt.

2.3.2 Temperaturabhängigkeit von Widerständen

Die Werkstoffe, aus denen die Widerstände hergestellt werden, ändern im Allgemeinen ihren spezifischen Widerstand in Abhängigkeit von der Temperatur. Diese Temperaturabhängigkeit lässt sich durch das folgende Polynom allgemein beschreiben:

$$R = R_0 \cdot (1 + \alpha \cdot (\vartheta - \vartheta_0) + \beta \cdot (\vartheta - \vartheta_0)^2 + \gamma \cdot (\vartheta - \vartheta_0)^3 + \dots) \quad (2.4)$$

Dabei ist R_0 der Widerstandswert bei $\vartheta_0 = 20^\circ\text{C}$, ϑ die Temperatur, die der Widerstand beim Betrieb erreicht und α , β und γ Temperaturkoeffizienten, die wiederum selbst von der Temperatur abhängen können, also real keine echten Konstanten sind. Für übliche ohm'sche Widerstände können die Koeffizienten β , γ und höherer Ordnung vernachlässigt werden, womit sich für die Temperaturabhängigkeit folgende Formel ergibt:

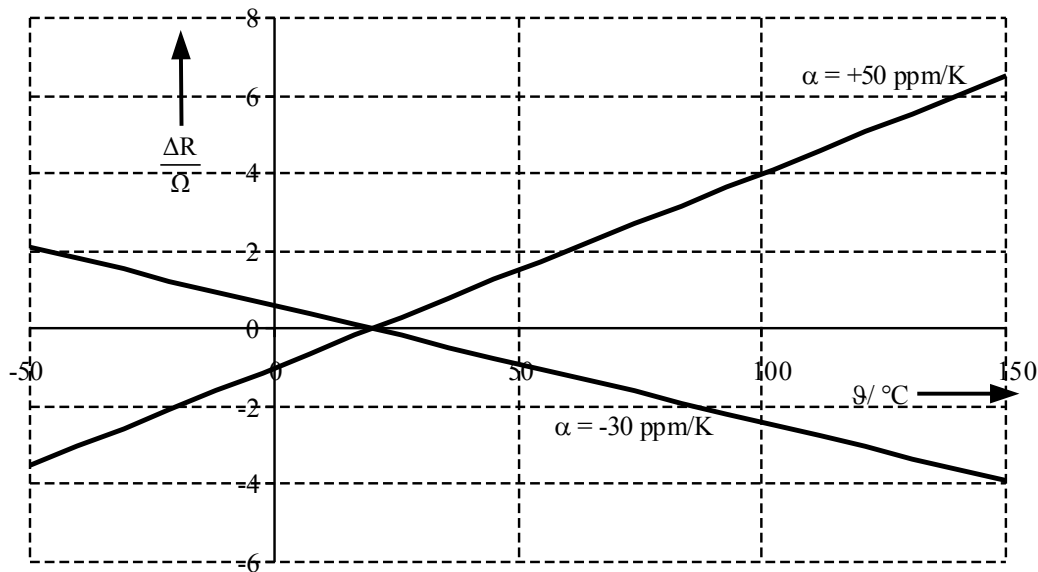


Bild 2.6: Temperaturabhängigkeit von zwei 1 kΩ – Widerständen mit verschiedenem Temperaturkoeffizienten α .

$$R = R_0 \cdot (1 + \alpha \cdot (\vartheta - \vartheta_0)) \quad \text{mit} \quad \alpha = \frac{\Delta R}{R} \cdot \frac{1}{\Delta \vartheta} \quad (2.5)$$

Der Temperaturkoeffizient α hat die Dimension $\frac{1}{K}$ und kann abhängig vom verwendeten Material positive oder negative Werte annehmen. In Bild 2.6 sind die temperaturabhängigen Änderungen der Widerstandswerte für zwei 1 kΩ–Widerstände mit a) $\alpha = -30 \text{ ppm/K}$ und b) $\alpha = +50 \text{ ppm/K}$ für einen Temperaturbereich von -50°C bis $+150^\circ\text{C}$ aufgetragen.

2.3.3 Belastbarkeit von Widerständen

Beim Betrieb von Widerständen wird die elektrische Leistung im Bauelement in Wärme umgewandelt. Dadurch erwärmt sich das Bauelement und wird thermisch belastet. Die im Widerstand umgewandelte Leistung wird auch als Verlustleistung bezeichnet, da sie in elektronischen Schaltungen nicht als Wärme genutzt wird (Ausnahmen: Heizofen, Herdplatte o.ä.). In den Datenblättern der Hersteller sind die maximal zulässigen Verlustleistungen für die einzelnen Typen und Bauformen von Widerständen angegeben, damit der Schaltungsentwickler verhindern kann, dass das Bauelement thermisch überlastet und damit beschädigt wird. Die erzeugte Wärme muss über die Oberfläche und die Anschlüsse an die Umgebung abgeführt werden. Zu beachten ist hierbei auch, dass die maximale Verlustleistung die im Widerstand erzeugt werden darf, von der Differenz der Oberflächentemperatur und der Umgebungstemperatur abhängt. Je geringer diese

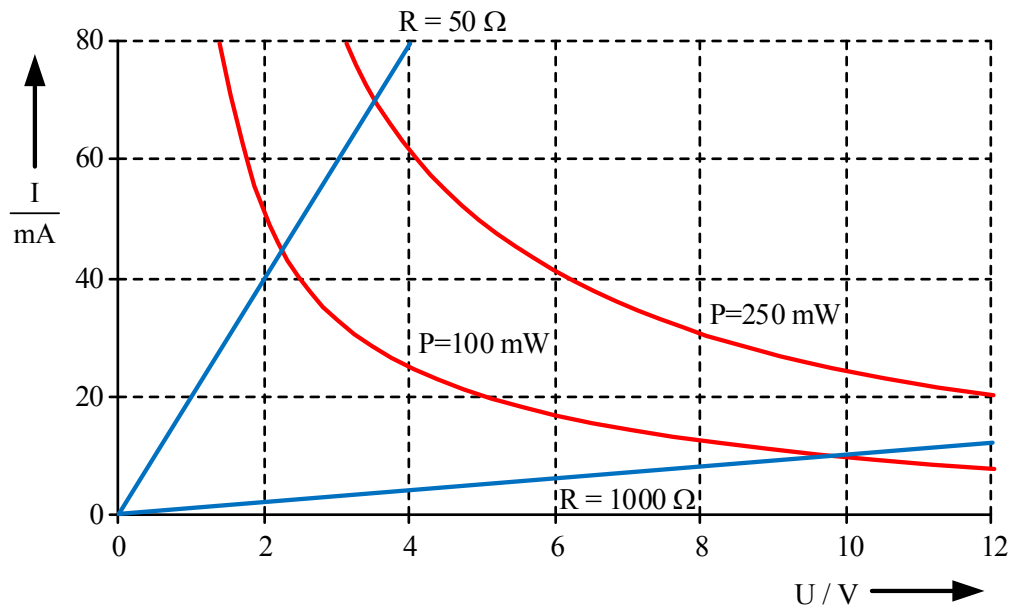


Bild 2.7: Strom-Spannungsdiagramm mit Verlustleistungshyperbeln für 100 mW und 250 mW.

Differenz ist, desto geringer ist auch die im Bauelement umgesetzte maximale Verlustleistung. Man spricht in diesem Fall auch von einem „Derating“ der Verlustleistung. Unter Verwendung von Gleichgrößen für Strom und Spannung erhält man:

$$\begin{aligned}
 P &= U \cdot I \quad \text{oder mit } U = R \cdot I \quad \text{wird} \quad P = R \cdot I^2 \\
 \text{bzw. mit } I &= \frac{U}{R} \quad \text{wird} \quad P = \frac{U^2}{R}
 \end{aligned}
 \tag{2.6}$$

Trägt man die Nennbelastung in ein Strom-Spannungsdiagramm ein, erhält man die Grenze des erlaubten Arbeitsbereiches. In Bild 2.7 sind je eine Kennlinie für einen Widerstandswert von $50 \, \Omega$ und $1000 \, \Omega$ und je eine Kurve konstanter Leistung, auch Verlustleistungshyperbel genannt, für 100 mW und 250 mW eingetragen.

2.3.4 Nichtlineare Widerstände

Neben den bisher beschriebenen linearen Widerständen werden für besondere Anwendungen auch Widerstände aus Werkstoffen hergestellt, die nichtlineare temperaturabhängige Eigenschaften besitzen. Der Temperaturkoeffizient ist wesentlich größer als bei ohm'schen Widerständen, woraus folgt, dass sich die Widerstandswerte dieser Bauelemente bei Veränderungen der Temperatur stark ändern können. Man unterscheidet zwischen Materialien mit negativen Temperaturkoeffizienten, sogenannten Heißleitern, oder NTC-Widerständen, und mit positiven Temperaturkoeffizienten, sogenannten Kaltlei-

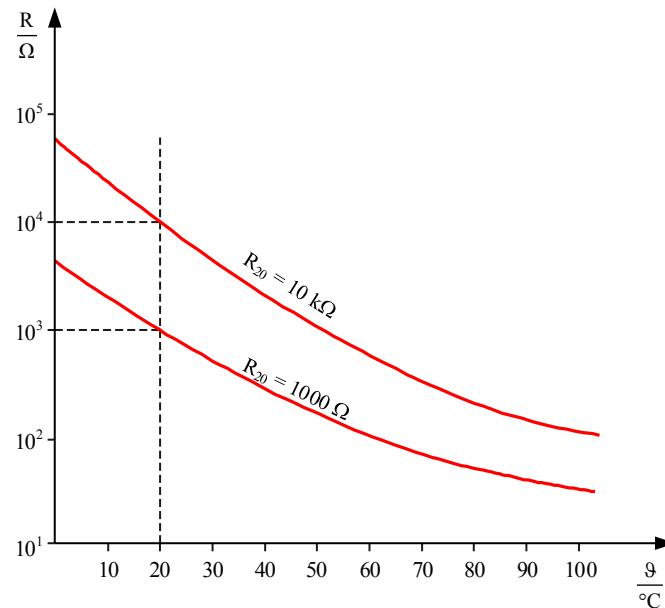


Bild 2.8: Widerstandsverläufe zweier NTC–Widerstände über der Temperatur.

tern oder PTC–Widerständen.

Ein Heißleiter hat bei hohen Temperaturen einen besonders kleinen Widerstand. Bild 2.8 zeigt die Widerstandsverläufe zweier NTC–Widerstände. Typische Werte von α liegen in der Größenordnung von -2 ‰/K bis -10 ‰/K . Heißleiterwiderstände werden in Stromkreisen zur Herabsetzung des Einschaltstroms, als Temperaturfühler und zur Temperaturstabilisierung von Halbleiterschaltungen eingesetzt. Kaltleiter oder PTC–Widerstände dagegen leiten bei niedrigen Temperaturen besonders gut, d.h. bei niedrigen Temperaturen ist ihr Widerstandswert besonders niedrig. Bild 2.9 zeigt eine Kennlinie eines Kaltleiters. Der Widerstandwert kann sich bei Erwärmung im zulässigen Arbeitsbereich sogar um mehrere Größenordnungen verändern.

Der Verlauf der Widerstandsänderung über der Temperatur ist stark nichtlinear. Während der Widerstandswert bei Temperaturen wenig über 20°C zunächst sogar leicht absinkt steigt er dann bei weiterer Erhöhung der Temperatur schnell um mehrere Größenordnungen an. Typische Werte von α liegen in der Größenordnung von $+7\text{ ‰/K}$ bis $+50\text{ ‰/K}$.

PTC–Widerstände werden auf zwei Arten betrieben: Eigenerwärmung und Fremderwärmung. Eine Eigenerwärmung wird erreicht, in dem man an das Bauelement Spannungen im Bereich zwischen 10 V und 60 V anlegt, die so groß sind, dass sich das Bauelement durch den Strom merklich selbst erwärmt, und so seinen Widerstandwert ändert. Bei Fremderwärmung müssen Spannung und Strom hingegen so klein gewählt werden, dass eine Eigenerwärmung ausgeschlossen ist. Auf diese Weise kann ein Kalt-

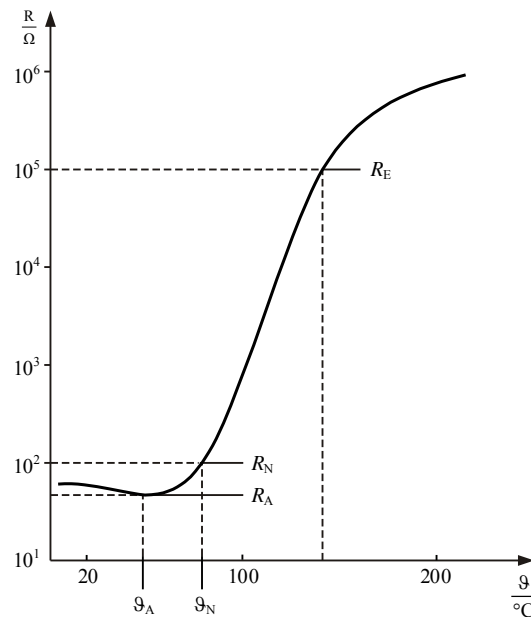


Bild 2.9: Widerstandsverlauf eines PTC–Widerstandes über der Temperatur.

leiter als Temperaturfühler eingesetzt werden.

2.4 Kondensatoren

Wie ohm'sche Widerstände haben auch Kondensatoren in der Regel zwei Anschlüsse. Sie bestehen aus zwei leitfähigen Platten bzw. Ebenen mit der Fläche A , die sich in einem Abstand d gegenüber liegen und voneinander isoliert sind. Die einfachste Bauform ist der Plattenkondensator nach Bild 2.10. Kondensatoren können elektrische Ladungen speichern. Es gilt:

$$\partial Q = i \cdot \partial t = C \cdot \partial u \quad \text{und} \quad C = \varepsilon_0 \cdot \varepsilon_r \cdot \frac{A}{d} \quad (2.7)$$

C hat die Dimension: $[C] = \frac{[i][dt]}{du} = \frac{As}{V} = F$

Die Kapazität C wird in Farad angegeben. Die Abhängigkeiten von zeitlich veränderlichen Strömen und Spannungen am Kondensator lassen sich folgendermaßen angeben:

$$i = C \cdot \frac{\partial u}{\partial t} \quad \text{und} \quad u = \frac{1}{C} \cdot \int_0^t i \cdot dt \quad (2.8)$$

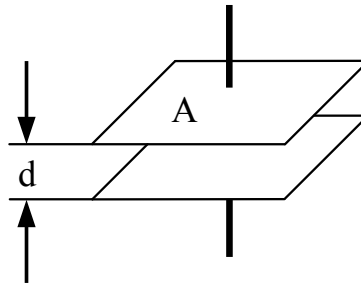


Bild 2.10: Schematischer Aufbau eines Plattenkondensators.

Schaltungstechnisch ist beim Einsatz von sinusförmigen Strömen bzw. Spannungen der Blindwiderstand des Kondensators durch

$$X_C = \frac{1}{\omega \cdot C} \quad (2.9)$$

gegeben. Ist die Frequenz sehr niedrig, wird X_C sehr groß und geht bei $f \rightarrow 0$ gegen unendlich große Werte ($X_C \rightarrow \infty$). Dagegen kann der Blindwiderstand eines Kondensators bei sehr hohen Frequenzen ($f \rightarrow \infty$) näherungsweise oft als Kurzschluss mit $X_C \rightarrow 0$ betrachtet werden.

Große Kondensatoren werden zur Speicherung von Ladung in Netzgeräten zur Stromversorgung von elektronischen Schaltungen und auf Platinen zur Pufferung der Versorgungsspannung der integrierten Bauelemente eingesetzt. Gleichung 2.8 zeigt, dass ein Strom auf die Platten eines Kondensators nur dann fließt, wenn eine zeitlich veränderliche Spannung angelegt wird, das heißt der Kondensator wirkt differenzierend. Wird ein Kondensator jedoch mit einem konstanten Strom über eine bestimmte Zeit t aufgeladen, so ist die zu messende Spannung gleich dem Integral des Stromes i über der Zeit t . Der Kondensator wirkt jetzt wie ein analoger Integrator. Beide Verhaltensweise können anhand der Netzwerke in Bild 2.11 verdeutlicht werden.

2.4.1 Bauformen von Kondensatoren

Kondensatoren werden je nach Kapazitätsbereich und Anwendungszweck mit unterschiedlichen Materialien hergestellt. Werden hohe Kapazitätswerte für niedrige bis mittlere Spannungsfestigkeit benötigt, verwendet man Elektrolytkondensatoren. Ist jedoch eine hohe Spannungsfestigkeit bis in den kV-Bereich erforderlich, setzt man Metall-Papier und Folienkondensatoren ein. Im Bereich der Elektronikschaltungen bei niedrigen Spannungen verwendet man Tantal-Elektrolyt-, Folien- oder Keramikkondensatoren,

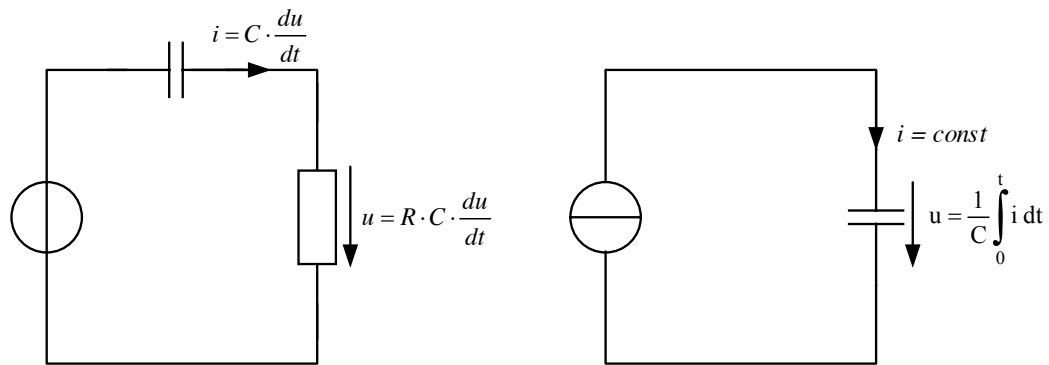


Bild 2.11: Kondensator als Differentiator und Integrator.

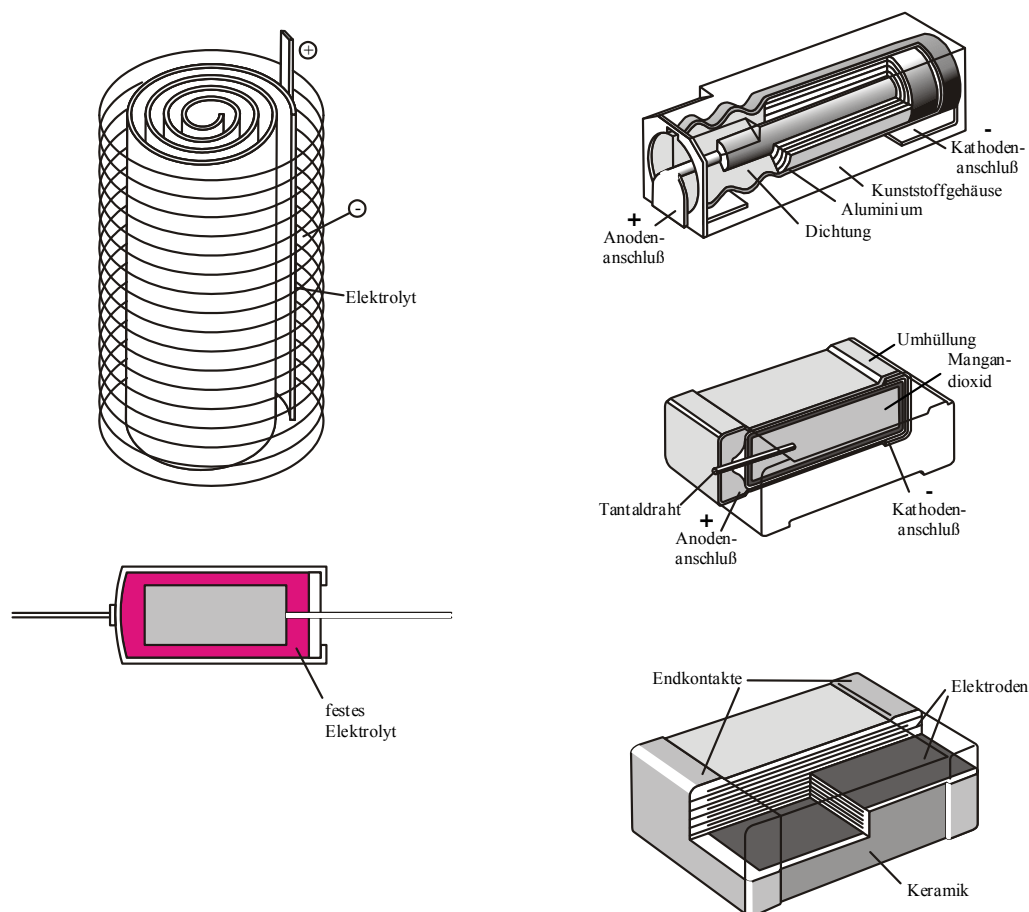


Bild 2.12: Bauformen verschiedener Kondensatoren.

abhängig vom benötigten Kapazitätswert und dem Einsatzgebiet. Bild 2.12 zeigt verschiedene Bauformen von Kondensatoren.

2.5 Spulen

Als letztes Bauelement aus der Gruppe der passiven Bauelemente soll die Spule bzw. Induktivität besprochen werden. Auch sie hat im Allgemeinen zwei Anschlüsse. Induktivitäten werden überwiegend aus mit Lack isoliertem Kupferdraht hergestellt und um einen Spulenkörper gewickelt. Die Induktivität der Spule ist abhängig von der Anzahl der Windungen n , der mittleren Länge ℓ , der Querschnittsfläche A des Spulenkörpers und dem Material des Spulenkerns.

$$L = n^2 \cdot \mu_0 \cdot \mu_r \cdot \frac{A}{\ell} \quad (2.10)$$

Wird eine Spule von einem sich zeitlich ändernden Strom durchflossen, so entsteht in ihrer Umgebung ein sich zeitlich änderndes Magnetfeld. Dieses magnetische Feld induziert in die Spule eine Spannung. Die Größe der induzierten Spannung ergibt sich aus dem Induktionsgesetz:

$$u_i = -n \cdot \frac{\partial \Phi}{\partial t} \quad (2.11)$$

Die induzierte Spannung u_i ist also der Änderung ihrer Ursache entgegen gerichtet. Unter Berücksichtigung des zeitlich veränderlichen Stromes i wird die induzierte Spannung zu

$$u_i = -L \cdot \frac{\partial i}{\partial t} \quad (2.12)$$

Die Einheit der Induktivität L ist Henry:

$$[L] = \frac{Vs}{A} = \Omega \cdot s = H \quad (2.13)$$

Der Blindwiderstand einer Spule ist durch die Beziehung

$$X_L = \omega \cdot L \quad (2.14)$$

definiert, d.h. für Gleichspannungen ist $X_L = 0$ und für sehr hohe Frequenzen geht $X_L \rightarrow \infty$.

2.5.1 Bauformen

In Bild 2.13 sind einige gebräuchliche Bauformen von Spulen dargestellt.

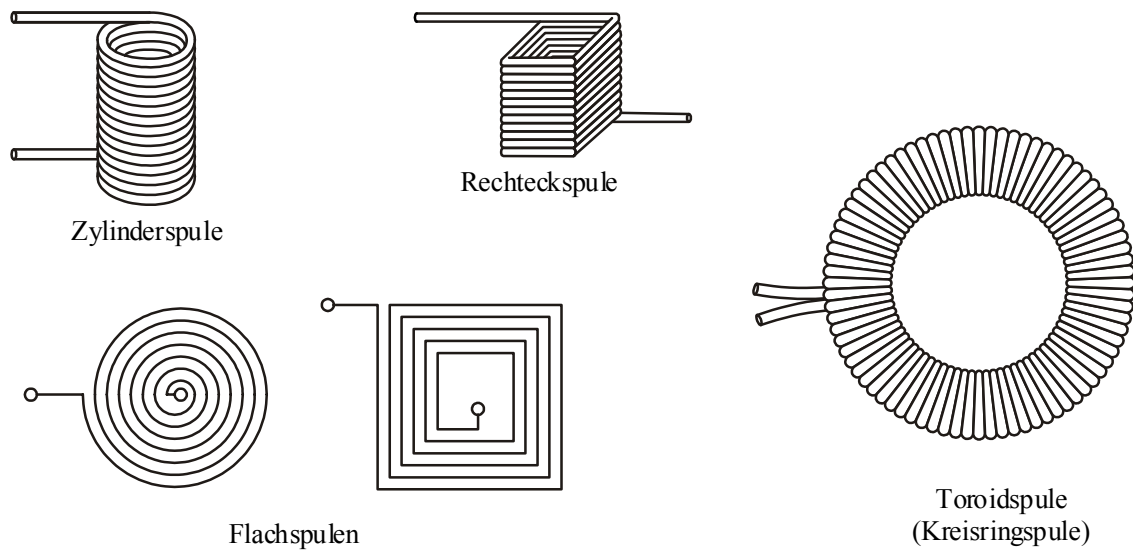


Bild 2.13: Bauformen von Spulen.

3 Halbleiterbauelemente

In der Entwicklung von integrierten Schaltkreisen werden standardisierte Modelle für Dioden, Bipolar- und Feldeffekttransistoren verwendet. Die einzelnen Herstellungsprozesse unterscheiden sich dann nur in ihren Modellparametern, die in den Simulationsprogrammen benutzt werden. Diese Parameter bestimmt der Halbleiter-Hersteller aus dem laufenden Herstellungsprozess und stellt sie der eigenen Entwicklung von integrierten Schaltkreisen zur Verfügung. In dem vorliegenden Kapitel gehen wir schrittweise vor, indem wir ausgehend von der Beschreibung und einfachen Modellen zur Funktion der Bauelemente deren phänomenologische Grundlage und Funktion beschreiben und dann zu den Ersatzschaltbildern überleiten, wie sie zumindest als Basisversion im üblichen Simulationsprogramm, wie PSPICE, benutzt werden.

Im Rahmen dieser Vorlesung verzichten wir auf die detaillierte Behandlung der festkörperphysikalischen Grundlagen des Aufbaus und der Funktion der einzelnen Bauelemente und einer Diskussion des Zusammenhangs der Modellparameter mit diesen physikalischen Grundlagen. Diese Fragen werden in der Vorlesung „Halbleiter-Bauelemente“ behandelt.

Einige physikalische Eigenschaften und Grundlagen der Halbleiterphysik sollen jedoch an dieser Stelle erwähnt werden, da diese das Verhalten von Halbleiter-Bauelementen bei ihrer Anwendung in elektronischen Schaltungen beeinflussen und deshalb berücksichtigt werden müssen. Die Leitfähigkeit des Halbleitermaterials ist der Konzentration der darin befindlichen Leitungsträger proportional. Damit wird es möglich, über Dotierung, also den Einbau von Fremdatomen, die Leitfähigkeit zu beeinflussen. Wir beschränken uns im vorliegenden Fall auf den Halbleiter Silizium (Si). Zur Herstellung von Halbleiter-Bauelementen wird das reine Halbleitermaterial Si mit fremden Atomen dotiert. Damit entstehen sog. Majoritäts- und Minoritäts-Ladungsträger. Silizium ist 4-wertig und durch die Zugabe eines 5-wertigen Donatorelementes, z.B. Phosphor, wird ein Elektron des Donators freigesetzt und deshalb wird das Silizium n-leitend. Gibt man im Gegensatz dazu dem Silizium einen sog. Akzeptor, z.B. Bor, hinzu, der 3-wertig ist, entsteht ein sog. Loch, d.h. ein fehlendes Elektron, und damit entsteht ein p-leitendes Silizium. Deshalb sprechen wir auch im Falle des mit Bor dotierten Siliziums von p-Leitung und im Falle eines mit Phosphor dotierten Siliziums von n-Leitung. Die Temperaturabhängigkeit der Leitfähigkeit wird bestimmt durch die Konzentration der freien Ladungsträger und natürlich auch von ihrer Beweglichkeit. Im Silizium sind die freien Überschuss-Elektronen etwa 2,5-mal beweglicher als die freien Löcher.

3.1 Die Diode

Lernziele

- Kennenlernen der wichtigsten Diodenarten und deren Eigenschaften
- Arbeitspunkteinstellung und Kleinsignalverhalten von Dioden
- Verstehen des Verhaltens klassischer Gleichrichterschaltungen
- Kennenlernen von Schaltungen zur Spannungsstabilisierung mit Zener-Dioden

3.1.1 Der Aufbau

Bringt man ein p- und ein n-leitendes Halbleitermaterial in Kontakt, entsteht an der Berührungsfläche durch einen Diffusionsprozess eine Übergangszone, in der sich die Elektronen und Löcher durch Diffusionsprozesse vermischen. Dieser Diffusionsprozess ist proportional zur Temperatur und hat eine charakteristische Energie von $k_B \cdot T \approx 26 \text{ meV}$ bei Zimmertemperatur ($T = 300 \text{ K}$). Die zugehörige Spannung $U_T = 26 \text{ mV}$ wird Temperaturspannung genannt. Im Verlauf des Diffusionsprozesses der Elektronen und Löcher bildet sich in der Übergangszone ein elektrisches Feld heraus. Mit zunehmender Trennung der Ladungsträger vergrößert sich das elektrische Feld, bis ein Gleichgewicht zwischen Elektronen und Löchern entsteht. Der Diffusionsprozess wird beendet, wenn folgende Bedingung erfüllt ist:

$$N_A \cdot x_p = N_D \cdot x_n$$

wobei N_A und N_D die Konzentration der Akzeptoren und Donatoren und x_n und x_p die Breite der entsprechenden Übergangszonen im Halbleiter darstellt. Wie beschrieben, entsteht ohne Anlegen einer äußeren Spannung damit ein elektrisches Feld, was zu einer Spannung über der Verarmungszone, die eine Raumladungszone darstellt, führt. Diese Spannung wird Diffusionsspannung genannt.

$$U_{Diff} = U_T \cdot \ln \left(\frac{N_A \cdot N_D}{n_i^2} \right)$$

Im Silizium beträgt diese Spannung etwa 0,6 bis 0,7 V und im Germanium z.B. nur etwa 0,3 V. In Bild 3.1 ist schematisch der pn-Übergang, die Ausbildung der Raumladungszone, die damit verbundene Spannung in Abhängigkeit vom Ort über dem pn-Übergang und das Schaltsymbol dargestellt. Die Diode ist ein Zweipolbauelement und die Anschlüsse werden mit Anode (A) und Kathode (K) bezeichnet. Eine Diode kann je nach angelegter positiver oder negativer Spannung entweder im Durchlass- oder im Sperrbetrieb betrieben werden. Im Sperrbetrieb vergrößert sich die Raumladungszone, was einer Ladungstrennung in dem Übergangsbereich entspricht. Diese Ladungstrennung

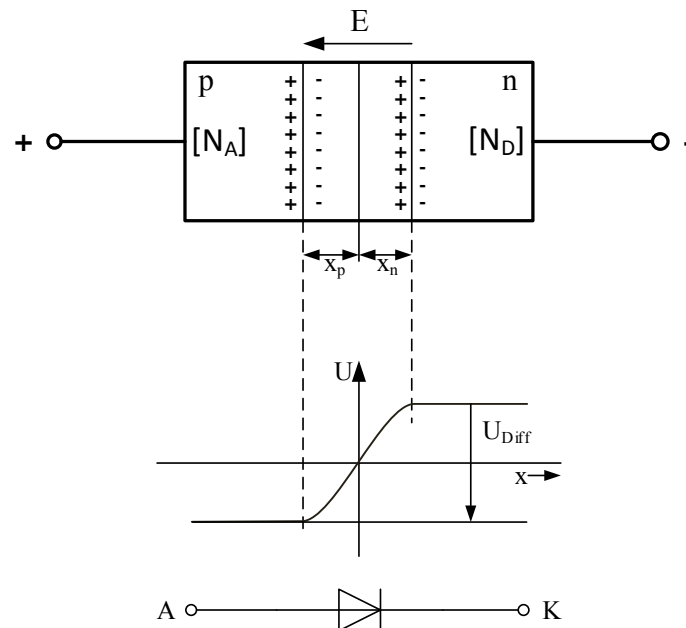


Bild 3.1: Schematische Darstellung der Raumladungszone, des Potentialverlaufes $U(x)$ und des Schaltsymbols eines pn-Überganges.

kann auch als Kapazität verstanden werden, was z.B. in Kapazitäts-Dioden ausgenutzt werden kann, um durch eine angelegte Spannung die Kapazität gezielt zu verändern. Ersetzt man den p-dotierten Teil des Halbleiters durch Aluminium erhält man eine Schottky-Diode, die sich durch eine sehr schmale Raumladungszone auszeichnet. In Bild 3.2 sind die beiden Übergänge schematisch dargestellt.

3.1.2 Die Kennlinien

Wird an eine Diode eine äußere Spannung angelegt, wird das innere elektrische Feld über der Raumladungszone beeinflusst. Durch das Anlegen einer Spannung in Durchlass-

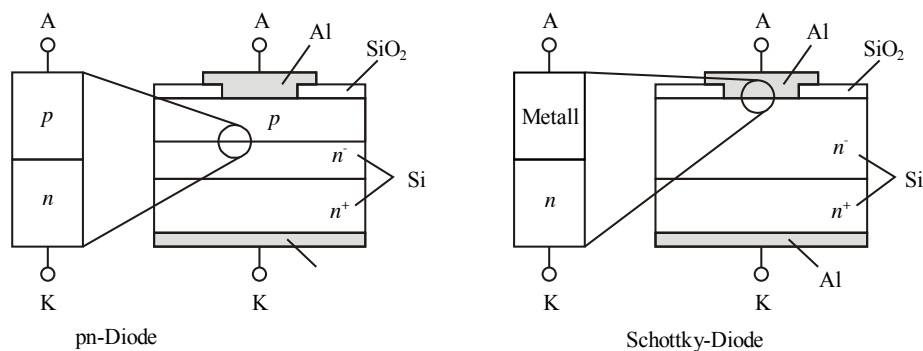


Bild 3.2: Prinzipieller Aufbau von pn- und Schottky-Dioden.

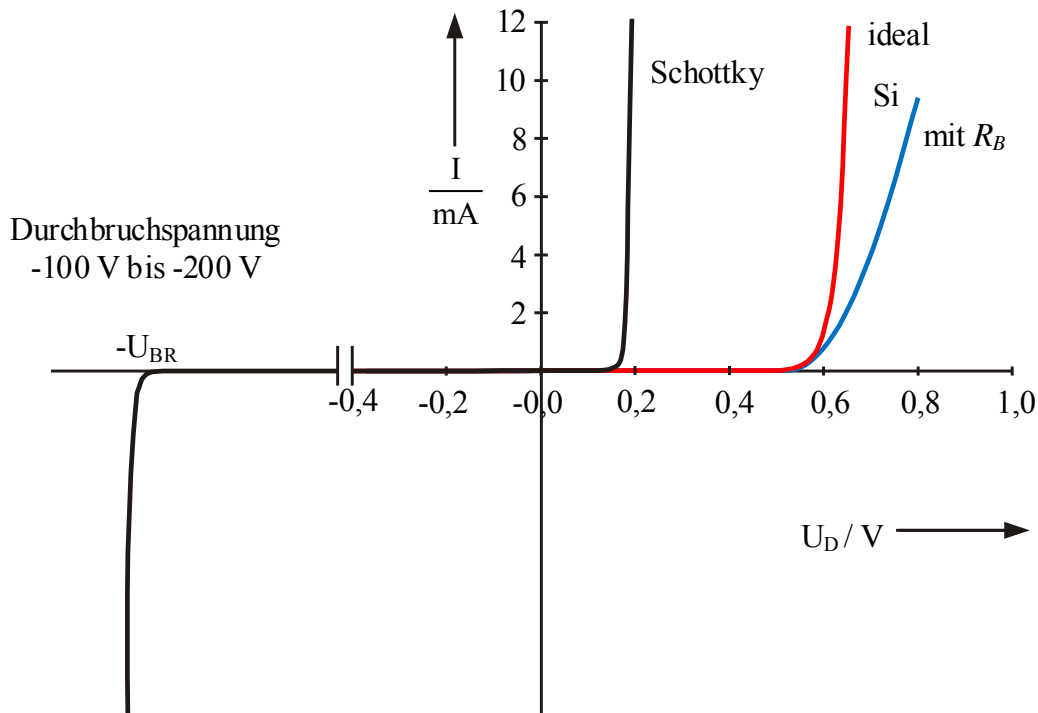


Bild 3.3: Strom-Spannungskennlinien einer pn-Diode und einer Schottky-Diode.

richtung (Anode +, Kathode –), wirkt diese äußere Spannung der Diffusionsspannung entgegen und die Breite der Raumladungszone wird kleiner. Dadurch kann ein Strom I_D durch den pn-Übergang fließen, der im Idealfall exponentiell mit der angelegten Spannung ansteigt (Gl. 3.1). Nimmt man jedoch an, dass das Halbleitermaterial einen endlichen, wenn auch kleinen, ohm'schen Bahnwiderstand R_B besitzt, wirkt sich dieser bereits ab etwa 0,5 V auf die Form der Kennlinie aus, sie wird etwas flacher. Polt man die angelegte Spannung um (Anode –, Kathode +) wird die Wirkung des inneren elektrischen Feldes noch verstärkt und die Breite der Raumladungszone wird größer. Es fließt jetzt nur ein sehr kleiner Sperrstrom I_S durch die Diode. Dieser Arbeitsbereich wird als Sperrbereich bezeichnet. In Bild 3.3 sind die Kennlinien einer Schottky-Diode und einer Si pn-Diode dargestellt. Eine Spannungen $U_D < -U_{BR}$ führt ebenfalls zu einem extrem schnellen Anstieg des negativen Stromes durch die Diode. Deshalb wird dieser Bereich als Durchbruchbereich bezeichnet. Liegt $-U_{BR}$ im Bereich von wenigen Volt wird dies für spezielle Dioden (Zener-Dioden) ausgenutzt. Dann wird dieser Bereich Zener-Bereich genannt. Trägt man die Kennlinie halblogarithmisch auf, so erhält man für den Bereich $U_D > 0$ annähernd eine Gerade, wie es in Bild 3.4 dargestellt ist. Unter Berücksichtigung des Bahnwiderstandes R_B zeigt die Kennlinie bei etwa $20 \cdot U_T$ ein Abweichen vom exponentiellen Anstieg. Im Sperrbereich, wird schnell der Sättigungswert des Sperrstroms I_S erreicht.

$$I = I_S \cdot \left(e^{\frac{U}{n \cdot U_T}} - 1 \right) \quad \text{mit} \quad U_T = \frac{k_B \cdot T}{e} \quad (3.1)$$

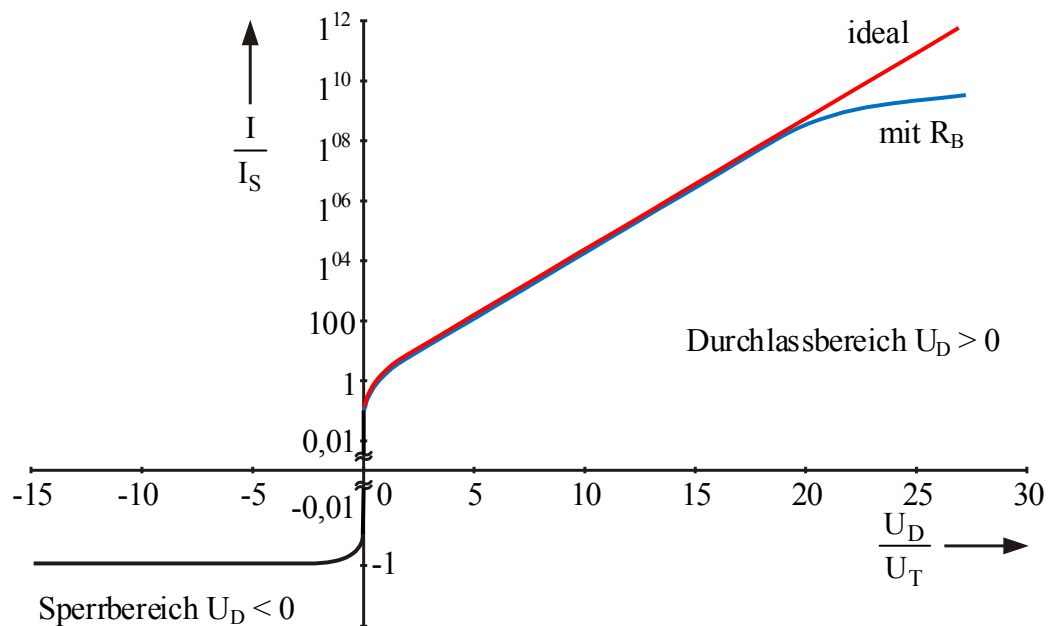


Bild 3.4: Diodenkennlinie im halblogarithmischen Maßstab.

Dabei ist $I_S \cong 10^{-12} \dots 10^{-6}$ A der Sättigungssperrstrom, $n \cong 1 \dots 2$ der Emissionskoeffizient und U_T die Temperaturspannung. Im Sperrbereich für Spannungen $U_D < 0$ fließt ein konstanter, kleiner Sperrstrom. Im Durchlassbereich für Spannungen U_D größer als $n \cdot U_T \cong 26 - 52$ mV gilt folgende Näherung:

$$I = I_S \cdot \left(e^{\frac{U}{U_T}} - 1 \right) \quad (3.2)$$

Wie bereits erwähnt, ist die Leitfähigkeit des Halbleitermaterials stark abhängig von der Temperatur. Aus diesem Grund wird oft der sog. Temperaturkoeffizient bei konstanter Spannung bzw. der Temperaturkoeffizient bei konstantem Strom einer Diode bestimmt. Der Temperaturkoeffizient des Stromes bei konstanter Spannung beträgt:

$$\frac{\partial I}{I \partial t} = \left(\frac{1}{I} \cdot \frac{\Delta I}{\Delta T} \right) \Big|_{U=\text{const.}} = 0,075/\text{K} = 7,5\%/\text{K}$$

Der Einfluss der Temperatur auf die Spannung einer Diode bei konstantem Strom beträgt:

$$\frac{\partial U}{\partial t} = \left(\frac{\Delta U}{\Delta T} \right) \Big|_{I=\text{const.}} = -2 \text{ mV/K}$$

d.h. eine Temperaturerhöhung um 10 K führt dazu, dass sich der Strom einer Diode bei einer konstanten Spannung fast verdoppelt.

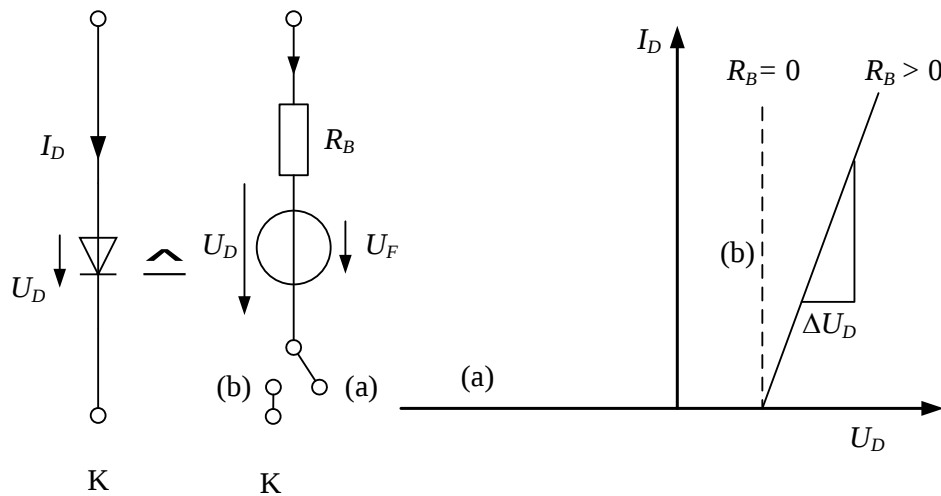


Bild 3.5: Vereinfachtes Ersatzschaltbild und zugehörige Strom-Spannungskennlinie einer Diode.

3.1.3 Ersatzschaltbild und statisches Verhalten

Für einfache Berechnungen kann die Diode als Schalter betrachtet werden, der im Sperrbereich geöffnet und im Durchlassbereich geschlossen ist. Nimmt man an, dass im Durchlassbereich die Spannung näherungsweise konstant ist und im Sperrbereich kein Strom fließt, kann man die Diode durch einen idealen spannungsgesteuerten Schalter und eine Spannungsquelle mit der Flussspannung U_F ersetzen (siehe Bild 3.5a). Das Bild 3.5b zeigt die Kennlinie dieser Ersatzschaltung, die sich aus zwei Geradenstücken zusammensetzt. Damit ergibt sich:

$$\begin{aligned} I_D &= 0 & \text{für } U < U_F \\ I_D &\rightarrow \infty & \text{für } U = U_F \end{aligned}$$

Die zugehörige Schaltung und die Kennlinie sind ebenfalls in Bild 3.5 gestrichelt dargestellt. Berücksichtigt man den Bahnwiderstand R_B ergibt sich folgende lineare Näherung:

$$I_D = \begin{cases} 0 & \text{für } U < U_F & \text{Schalter in Stellung (a)} \\ \frac{U - U_F}{R_B} & \text{für } U > U_F & \text{Schalter in Stellung (b)} \end{cases}$$

In beiden Varianten ist eine Fallunterscheidung notwendig, d.h. man muss mit offenem und geschlossenem Schalter rechnen und somit den entsprechenden Fall ermitteln, der nicht zu einem Widerspruch der Funktion der Schaltung führt. Der Vorteil dieses Ersatzschaltbildes besteht darin, dass mit einer linearen Gleichung gerechnet werden kann und damit einfache mathematische Lösungen und Simulationen möglich sind.

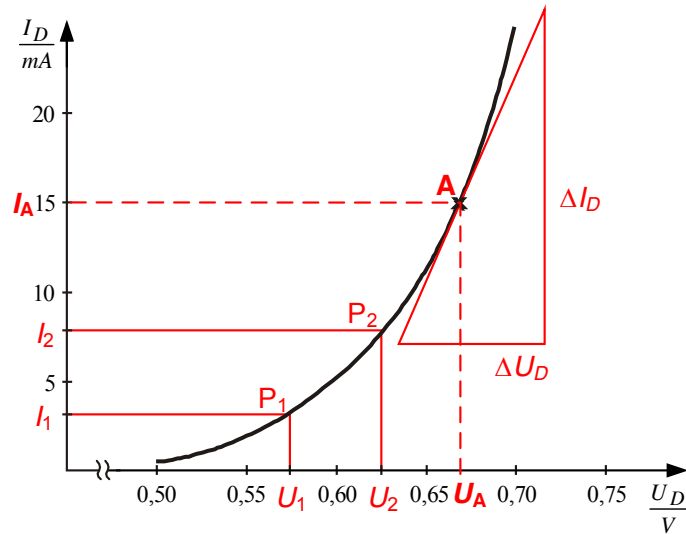


Bild 3.6: Linearisierung einer Diodenkennlinie im Arbeitspunkt A.

3.1.4 Das Kleinsignalverhalten

Die Kennlinie einer Diode ist in Bild 3.6 gezeigt. Der Arbeitspunkt der Diode wird durch eine feste Spannung und einen entsprechend zugeordneten Strom festgelegt. Wir bezeichnen den Arbeitspunkt mit dem Buchstaben A. Für den Einsatz in einer realen Schaltung interessiert in der Regel lediglich das Verhalten der Diode um diesen Arbeitspunkt A herum, d.h. bei Aussteuerung mit kleinen Spannungen und kleinen Strömen. Deshalb wird das Verhalten als Kleinsignalverhalten um diesen Arbeitspunkt herum bezeichnet. In diesem Fall kann die nichtlineare Kennlinie entsprechend Gleichung 3.1 durch eine Tangente im Arbeitspunkt A ersetzt werden. Dabei wird der differentieller Widerstand r_D wie folgt definiert:

$$r_D = \left. \frac{\partial U}{\partial I} \right|_A \cong \left. \frac{n \cdot U_T}{I_D} \right|_{I_D \gg I_S} \quad (3.3)$$

Demzufolge besteht das Kleinsignal-Ersatzschaltbild einer Diode nur aus dem differentiellen Widerstand r_D . Bei großen Strömen muss zu diesem Widerstand r_D noch zusätzlich der Bahnwiderstand R_B berücksichtigt werden. Damit kann man den Zusammenhang zwischen Spannung und Diodenstrom wie folgt schreiben:

$$U = I \cdot (r_D + R_B) \quad (3.4)$$

Da wir die bereits besprochene Sperrschicht-Kapazität der Raumladungszone nicht berücksichtigt haben, gilt dieses Kleinsignal-Ersatzschaltbild lediglich für niedrige Frequenzen bis ca. 10 kHz.

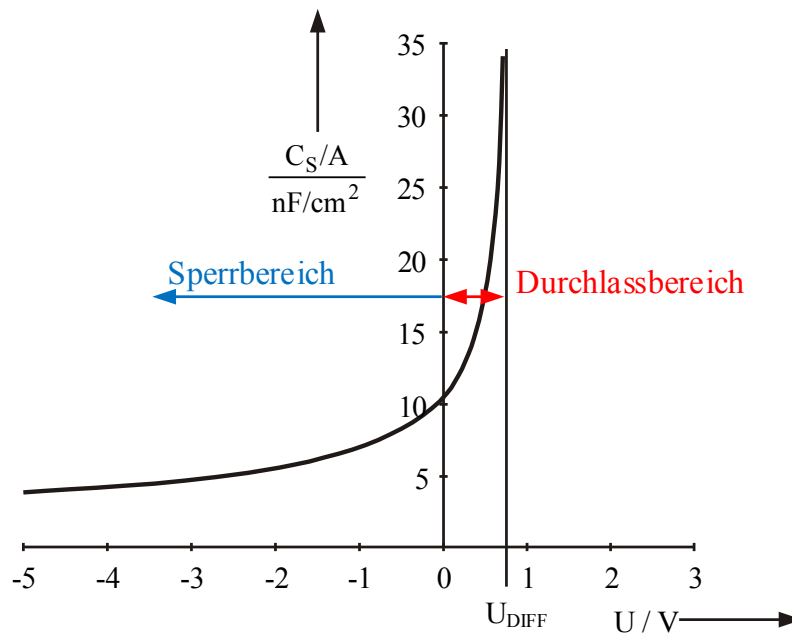


Bild 3.7: Abhängigkeit der Sperrschichtkapazität einer Diode von der Spannung.

3.1.5 Das dynamische Verhalten

Wie bereits besprochen wurde, sind die positiven und negativen Ladungen, also Löcher und Elektronen, in der Raumladungszone einer pn-Diode getrennt. Diese Raumladungszone hat damit eine Kapazität, weil eine bestimmte Ladungsmenge in diesem Volumen der Raumladungszone gespeichert ist. Diese Kapazität wird als Sperrschichtkapazität C_S bezeichnet. Sie lässt sich näherungsweise unter Annahme eines Parallelplattenkondensators wie folgt berechnen:

$$C = \varepsilon_r \cdot \varepsilon_0 \cdot \frac{A}{d} \quad (3.5)$$

Die Dicke der Raumladungszone d können wir genähert annehmen als $d \cong x_n + x_p$, was der Breite der Raumladungszone entspricht. Zusätzlich zur Sperrschicht-Kapazität fällt in einem pn-Übergang im Durchlassbereich eine Diffusionsladung auf, die proportional zum Diffusionsstrom durch den pn-Übergang ist. Die Diffusionskapazität C_D ist sehr viel kleiner als die Sperrschicht-Kapazität, muss aber im Bereich sehr hoher Frequenzen berücksichtigt werden. Bei Schottky-Dioden ist die Diffusionskapazität ca. 1000-mal kleiner als bei normalen pn-Dioden. Deshalb werden diese für Höchsthäufigkeiten eingesetzt. Bild 3.7 zeigt die Abhängigkeit der Sperrschichtkapazität von der Spannung an der Diode. Diese Abhängigkeit wird in so genannten Kapazitätsdioden ausgenutzt, wobei man durch eine angelegte Gleichspannung die Kapazität ändern und in Resonanzkreisen dazu benutzen kann, um z.B. die Resonanzfrequenz zu verstimmen.

3.1.6 Vollständiges Modell einer Diode

Das vollständige Modell einer Diode nach Bild 3.8 wird für die Schaltungssimulation mit Programmen wie z.B. PSPICE benötigt. Die Diodensymbole im Modell stehen für den Diffusionsstrom $I_{D,D}$ und den Rekombinationsstrom $I_{D,R}$; der Durchbruchstrom $I_{D,BR}$ ist durch eine gesteuerte Stromquelle dargestellt. In Tabelle 3.1 sind die Parameter des Diodenmodells aufgelistet, und die zusätzlichen Bezeichnungen der Parameter in der Schaltungssimulation PSPICE sind in Tabelle 3.2 angegeben. Tabelle 3.3 enthält Parameterwerte einiger ausgewählter Dioden, die der Bauteilbibliothek von PSPICE entnommen wurden. Nicht angegebene Werte in der Tabelle werden von PSPICE unterschiedlich behandelt, d.h. Werte 0 und ∞ bewirken, dass der jeweilige Effekt in der Simulation nicht berücksichtigt wird.

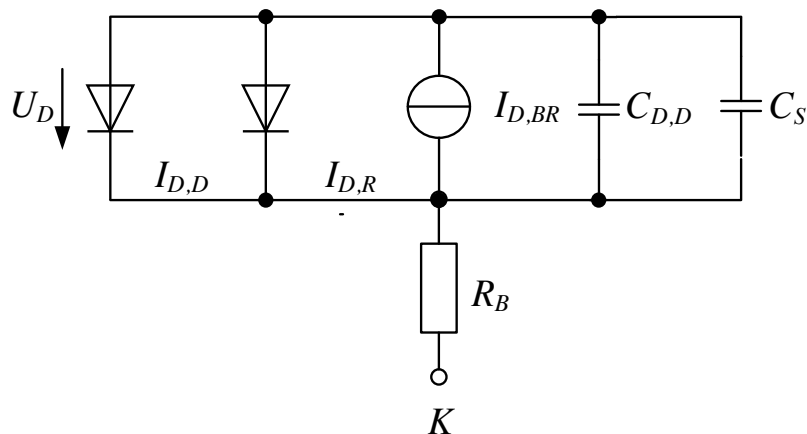


Bild 3.8: Vollständiges Ersatzschaltbild einer Diode.

Größe	Bezeichnung
$I_{D,D}$	Diffusionsstrom
$I_{D,R}$	Rekombinationsstrom
$I_{D,BR}$	Durchbruchstrom
R_B	Bahnwiderstand
C_S	Sperrschichtkapazität
$C_{D,D}$	Diffusionskapazität

Tabelle 3.1: Größen des Dioden-Modells.

Parameter	PSpice	Bezeichnung
Statisches Verhalten		
I_S	IS	Sättigungsstrom
n	N	Emissionskoeffizient
$I_{S,R}$	ISR	Leck-Sättigungsstrom
n_R	NR	Emissionskoeffizient
I_K	IK	Kniestrom zur starken Injektion
I_{BR}	IBV	Durchbruch-Kniestrom
n_{BR}	NBV	Emissionskoeffizient
U_{BR}	BV	Durchbruchspannung
R_B	RS	Banwiderstand
Dynamisches Verhalten		
C_{S0}	CJO	Null-Kapazität der Sperrschicht
U_{Diff}	VJ	Diffusionsspannung
m_s	M	Kapazitätskoeffizient
f_S	FC	Koeffizient für den Verlauf der Kapazität
τ_T	TT	Transitzeit
Thermisches Verhalten		
$x_{T,I}$	XTI	Temperaturkoeffizient der Sperrströme

Tabelle 3.2: Parameter des Dioden-Modells.

Parameter	PSpice	1N4148	1N4001	BAS40	Einheit
I_S	IS	2,68	14,1	0	nA
n	N	1,84	1,98	1	
$I_{S,R}$	ISR	1,57	0	254	fA
n_R	NR	2	2	2	
I_K	IK	0,041	94,8	0,01	A
I_{BR}	IBV	100	10	10	μ A
n_{BR}	NBV	1	1	1	
U_{BR}	BV	100	75	40	V
R_B	RS	0,6	0,034	0,1	Ω
C_{S0}	CJO	4	25,9	4	pF
U_{Diff}	VJ	0,5	0,325	0,5	V
m_S	M	0,333	0,44	0,333	
f_S	FC	0,5	0,5	0,5	
τ_T	TT	11,5	5700	0,025	ns
$x_{T,I}$	XTI	3	3	2	

Tabelle 3.3: Parameter einiger Dioden (1N4148: Kleinsignal-Diode, 1N4001: Gleichrichterdiode, BAS40: Schottky-Diode)

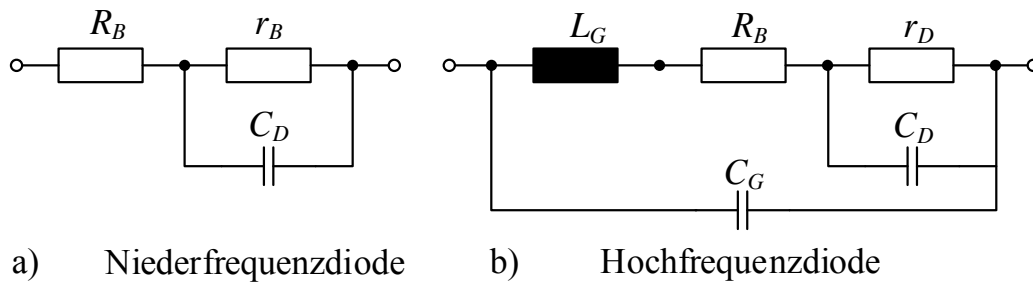


Bild 3.9: Dynamische Kleinsignalmodelle einer Diode.

3.1.7 Vereinfachtes dynamisches Kleinsignal-Modell

Wie bereits erwähnt, eignet sich das Kleinsignal-Modell für die meisten Simulationen, und deswegen soll es hier kurz eingeführt werden. Bild 3.9a zeigt das Kleinsignalmodell einer Niederfrequenz-Diode, in dem zusätzlich die Kapazität der Diode berücksichtigt wurde. In diesem Fall wurden entsprechend dem vollständigen Ersatzschaltbild nach Bild 3.8 die Sperrschichtkapazität und die Diffusionskapazität parallel geschaltet und mit C_D bezeichnet. Bei Hochfrequenzdioden muss man zusätzlich die parasitären Einflüsse des Gehäuses und der Zuleitungen berücksichtigen. Diese Effekte wurden in Bild 3.9b durch die Berücksichtigung der Gehäuseinduktivität L_G und der Gehäusekapazität C_G modelliert. Übliche Werte der Gehäuseinduktivität liegen bei etwa 1 – 10 nH und die der Gehäusekapazität bei etwa 0,1 – 1 pF. Für praktische Anwendungen reichen in der Regel folgende einfache Näherungen im Rahmen dieses dynamischen Kleinsignalmodells. Im Durchlassbereich beträgt die Diodenkapazität $C_D = 2 \cdot C_S$ und der differentielle Widerstand $r_D \cong n \cdot \frac{U_T}{I}$. Im Sperrbereich beträgt die Diodenkapazität $C_D \cong C_S$ und für den differentielle Widerstand gilt $r_d \rightarrow \infty$.

3.1.8 Anwendungsbeispiele

Die häufigste Anwendung einer Diode trifft man in sogenannten Gleichrichterschaltungen. Bild 3.10 zeigt die einfachste Gleichrichterschaltung, die Einweg-Gleichrichterschaltung. Die Wirkungsweise der Schaltung ist sehr einfach. Für positive Halbwellen einer Wechselspannung ist die Diode durchlässig und am Ausgang ist entsprechend positive Halbwelle zu sehen und für die negative Halbwelle sperrt die Diode und demzufolge ist im Ausgangssignal die negative Halbwelle nicht mehr enthalten. Den Spannungsverlauf der Ausgangsspannung sieht man ebenfalls in Bild 3.10. Der Nachteil dieser Einweg-Gleichrichterschaltung besteht darin, dass die Ausgangsspannung eine sehr große Welligkeit enthält. Um die Welligkeit der Ausgangsspannung zu verringern und gleichzeitig die negative Halbwelle mit auszunutzen, wird die sogenannte Brückenschaltung oder der Brückengleichrichter benutzt. Bild 3.11 zeigt die Brückengleichrichter-Schaltung, wie sie in eigentlich fast allen Netzteilen eingesetzt wird.

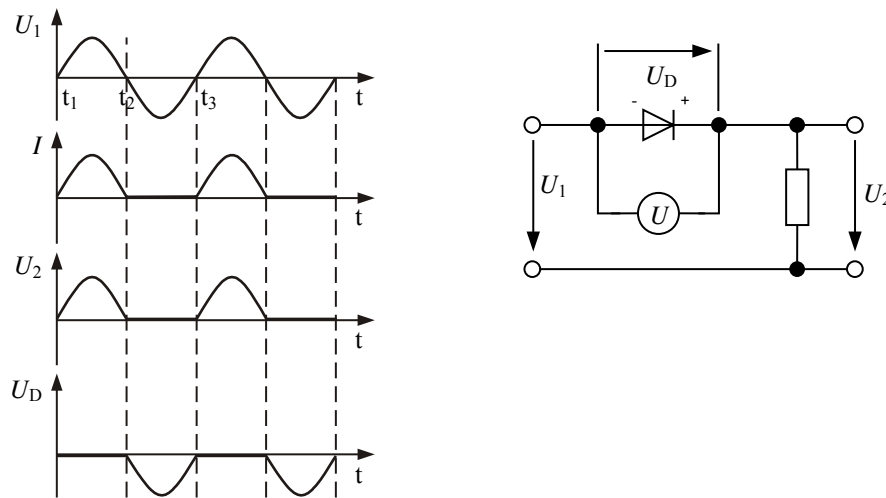


Bild 3.10: Einweg-Gleichrichterschaltung und zeitlicher Verlauf der Spannungen in der Schaltung. U_1 und U_2 bezeichnen entsprechend die Ein- und Ausgangsspannungen.

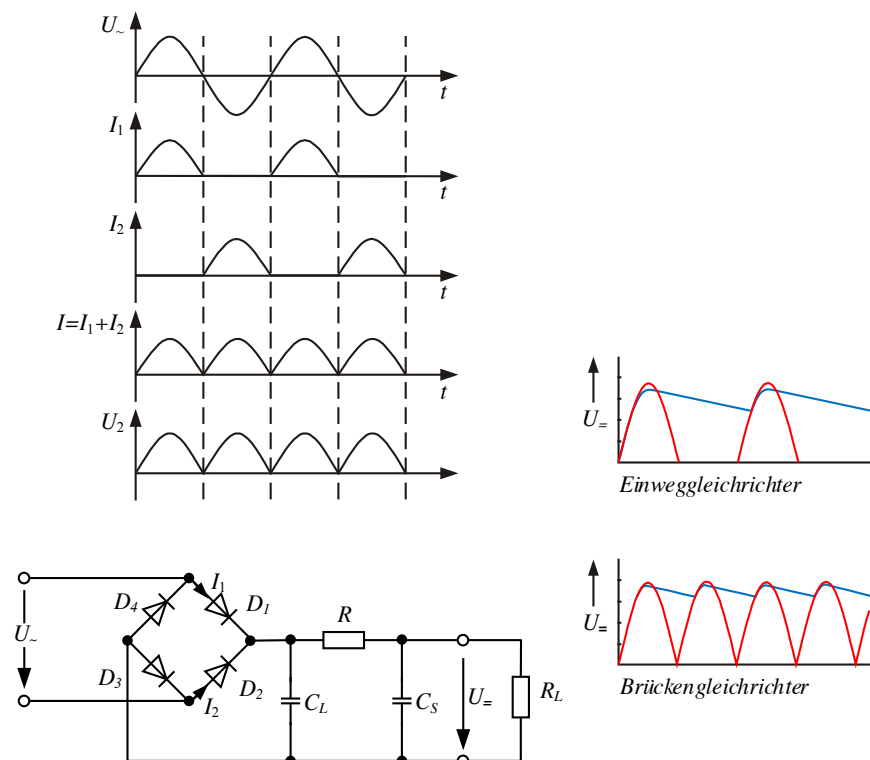


Bild 3.11: Brücken-Gleichrichterschaltung und zeitlicher Verlauf der Spannungen in der Schaltung. $U_{\sim}$ und $U_{=}$ bezeichnen entsprechend die Ein- und Ausgangsspannung.

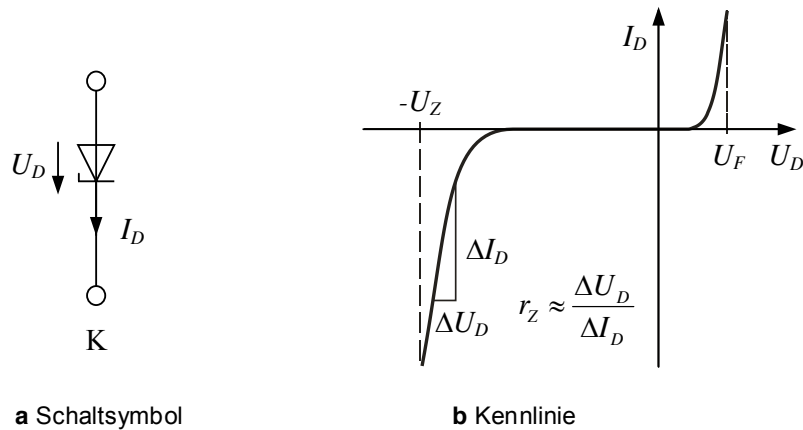


Bild 3.12: Schaltsymbol (a) und Kennlinie (b) einer Zener-Diode. U_Z ist die Zener-Spannung.

Bei positiven Eingangsspannungen leiten die Dioden D_1 und D_3 und bei negativen Halbwellen leiten die Dioden D_2 und D_4 . Die jeweils anderen Dioden sperren. Da der Strom immer über 2 leitende Dioden fließt, ist die gleichgerichtete Ausgangsspannung um den Betrag $2 \cdot U_f \cong 1,2$ bis 2 V kleiner als der Betrag der Eingangsspannung. Bild 3.11 zeigt ebenfalls den Verlauf der Ausgangsspannung nach der Brückengleichrichter-Schaltung. Man sieht, dass sich die Welligkeit im Mittleren wesentlich reduziert hat.

3.1.9 Die Zener-Diode

Zener-Dioden mit genau definierter Durchbruchspannung werden zur Spannungsstabilisierung bzw. zur Spannungsbegrenzung eingesetzt. Die Durchbruchspannung U_{BR} wird bei Zener-Dioden als Zener-Spannung U_Z bezeichnet und beträgt bei handelsüblichen Dioden 3 bis 30 V. Bild 3.12 zeigt das Schaltungssymbol und die Kennlinie einer Zener-Diode. Die Kennlinie einer Z-Diode kann mit folgender Gleichung im Durchbruchbereich beschrieben werden:

$$I_D \cong -I_{BR} \cdot e^{-\frac{U_D + U_Z}{n_{BR} \cdot U_T}} \quad (3.6)$$

Der differentielle Widerstand im Durchbruchbereich wird mit r_Z bezeichnet und lässt sich wie folgt darstellen:

$$r_Z = \frac{\partial U_D}{\partial I_D} = \frac{n_{BR} \cdot U_T}{I_D} \quad (3.7)$$

Der differentielle Widerstand hängt maßgeblich vom Emissionskoeffizienten n_{BR} ab. Für Zener-Dioden mit einer Zener-Spannung um etwa 3 V beträgt n_{BR} etwa 10 bis 20 und für Zener-Spannungen von ca. 47 V beträgt n_{BR} etwa 4 bis 8.

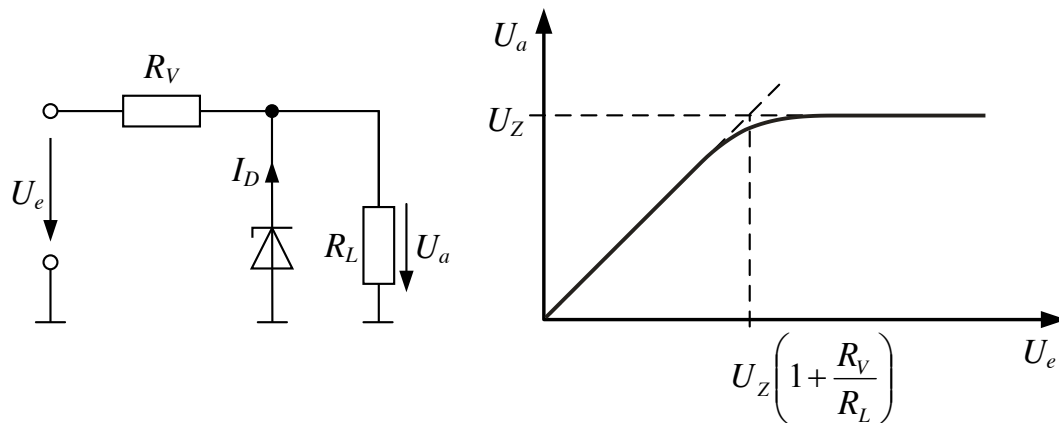


Bild 3.13: Schaltung zur Spannungsstabilisierung mit einer Z-Diode.

Bild 3.13 zeigt eine typische Schaltung zur Spannungsstabilisierung mit einer Zener-Diode. Für Eingangsspannungen kleiner als $U_Z > U_a \geq 0$ sperrt die Z-Diode und die Ausgangsspannung ergibt sich durch die Spannungsteilung an den Widerständen R_V und R_L . Wenn die Z-Diode leitet, gilt $U_a \cong U_Z$. Ein wichtiges Maß für die Beschreibung einer Schaltung zur Spannungsstabilisierung ist der Glättungsfaktor, der wie folgt definiert ist:

$$G = \frac{\Delta U_e}{\Delta U_a} \cong \frac{\partial U_e}{\partial U_a} \quad (3.8)$$

3.2 Bipolare Transistoren

Lernziele:

- Kennenlernen der wichtigsten Eigenschaften und Funktionen bipolarer Transistoren
- Verstehen und Arbeiten mit Kennlinien
- Einstellung des Arbeitspunktes und Kleinsignalverhalten
- Verstehen von Grenzdaten der Transistoren
- Grundschaltungen mit bipolaren Transistoren (Emitter-, Kollektor-, Basisschaltung)

3.2.1 Aufbau und Ersatzschaltbild

Bipolare Transistoren sind Halbleiterbauelemente mit den 3 Anschlüssen: Basis (B), Emitter (E) und Kollektor (C). Bipolare Transistoren kann man sich leicht als zwei gegenseitig geschaltete pn-Dioden vorstellen, die eine gemeinsame p- oder n-Zone besitzen. Bild 3.14 zeigt das Schaltzeichen und die Dioden-Ersatzschaltbilder von npn- und pnp-Transistoren. Dabei ist zu beachten, dass die Dioden-Ersatzschaltbilder nur die prinzipielle Funktion des Transistors wiedergeben, aber eine genaue Beschreibung durch dieses Schaltbild nicht erfolgt. Bipolare Transistoren werden zum Verstärken

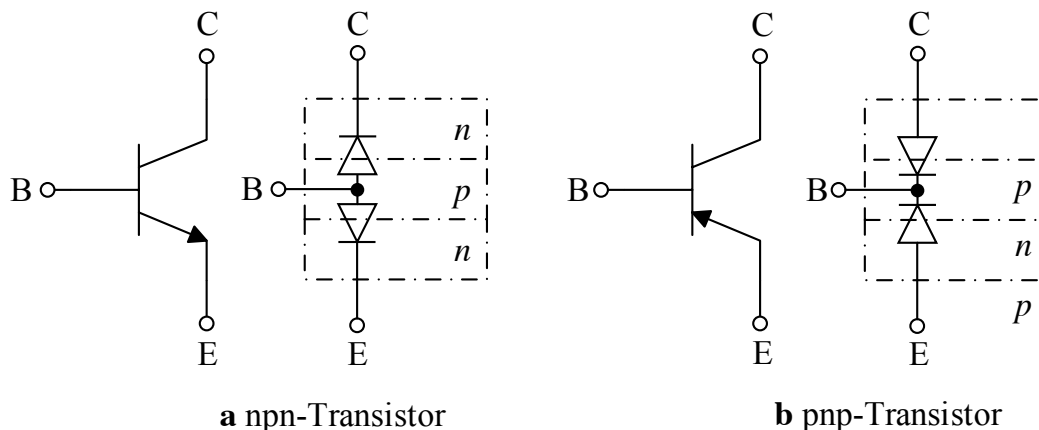


Bild 3.14: Dioden-Ersatzschaltbilder von n-p-n- und p-n-p - Transistoren.

und Schalten von Signalen eingesetzt und dabei meist im so genannten Normalbetrieb betrieben. Dabei ist die Emitterdiode in Flussrichtung und die Kollektordiode in Sperrrichtung geschaltet. Bild 3.15 zeigt die Richtung der Ströme für den Normalbetrieb von einem npn und einem pnp Transistor. Bild 3.16 zeigt den Aufbau von npn, pnp und integrierten Transistoren. Im Querschnitt sind die Bereiche mit n- und p-Leitung

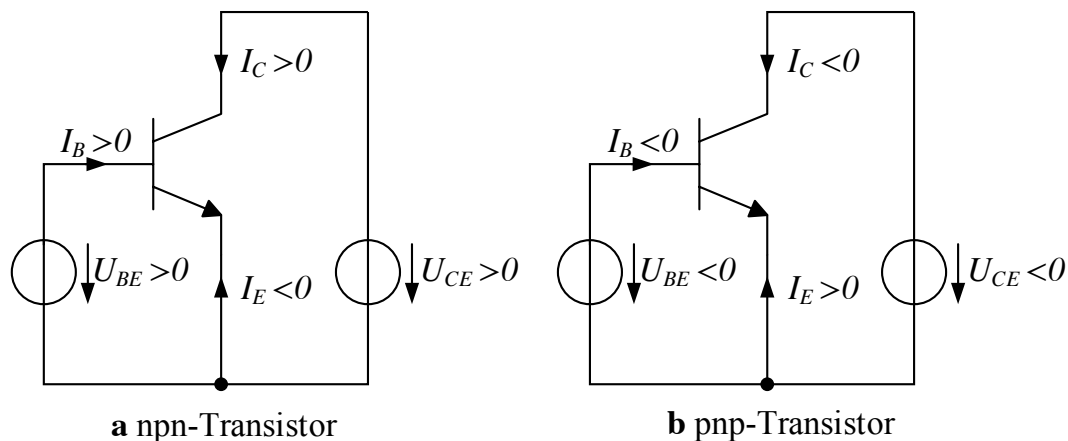


Bild 3.15: Spannungen und Ströme im Normalbetrieb.

besonders hervorgehoben; die Bereiche, die mit einem +, also $n+$ oder $p+$ bezeichnet sind, entsprechen einer sehr hohen Dotierung, um eine hohe Leitfähigkeit des Halbleitermaterials in diesem Bereich zu erreichen. Der Aufbau eines bipolaren Transistors in einem integrierten Schaltkreis erfordert zusätzliche Dotierungsschritte, um den Bereich des bipolaren Transistors von anderen Teilen der integrierten Schaltung zu trennen. Für diesen Zweck sind extra $p+$ Bereiche, die hier links und rechts markiert sind, angebracht und zur Kontaktierung spezielle $n+$ Bereiche hier zum Substrat speziell hervorgehoben. Die detaillierte Funktion der einzelnen Halbleiter-Dünnschichten und ihr Zusammenhang mit der Bauelementefunktion des bipolaren Transistors wird detailliert in der Vorlesung „Halbleiter-Bauelemente“ besprochen. Bild 3.17 zeigt verschiedene, weit verbreitete Gehäusebauformen von Einzeltransistoren.

3.2.2 Die Funktionen eines Bipolartransistors

3.2.2.1 Statische Kennlinien

Das Verhalten eines Bipolartransistors lässt sich am einfachsten anhand der Kennlinien aufzeigen. Sie beschreiben den Zusammenhang zwischen den Strömen und den Spannungen am Transistor für den Fall, dass alle Größen statisch, d.h. nicht oder nur sehr langsam zeitveränderlich sind. Legt man an einen Transistor entsprechend Bild 3.15a verschiedene Basis-Emitterspannungen U_{BE} an und misst jeweils den Kollektorstrom I_C als Funktion der Kollektor-Emitterspannung U_{CE} , so erhält man ein typisches Ausgangskennlinienfeld wie es in Bild 3.18 dargestellt ist. Mit der Ausnahme eines sehr kleinen Spannungsbereiches der I_C -Achse sind die Kennlinien nur wenig von U_{CE} abhängig und der Transistor arbeitet im Normalbetrieb, d.h. die Basis-Emitterdiode leitet und die Basis-Kollektordiode sperrt. Wenn die Kollektor-Emitterspannung U_{CE} den Wert der Basis-Emitterspannung U_{BE} unterschreitet, wird auch die Basis-Kollektordiode leitend und der Transistor gerät in die Sättigung. Bei der Sättigungsspannung $U_{CE,sat}$ knicken die Kennlinien scharf ab und verlaufen näherungsweise durch den Ursprung

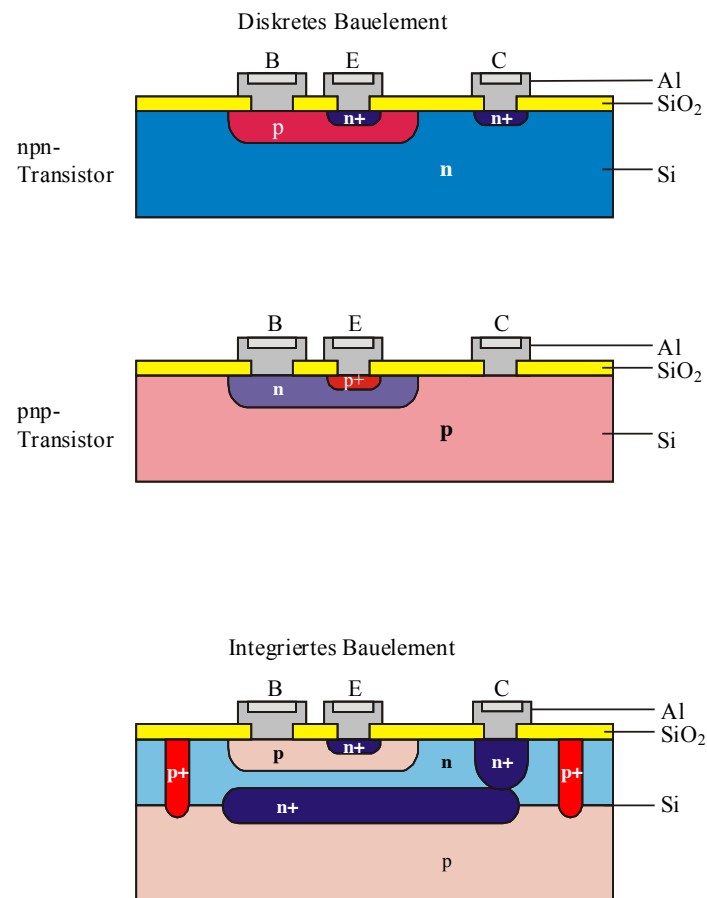


Bild 3.16: Schematische Ansicht des Aufbaus verschiedener bipolarer Transistoren.

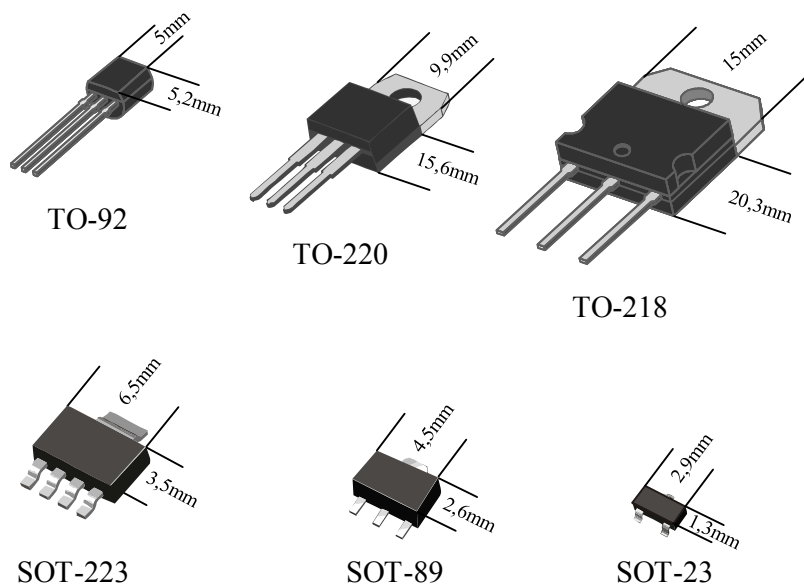


Bild 3.17: Bauformen diskreter Transistoren.

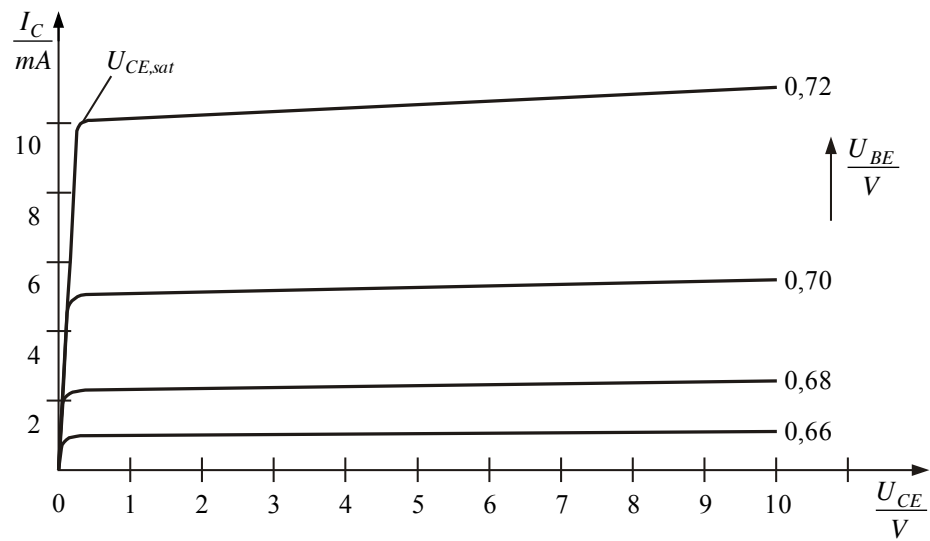


Bild 3.18: Spannungen und Ströme im Normalbetrieb.

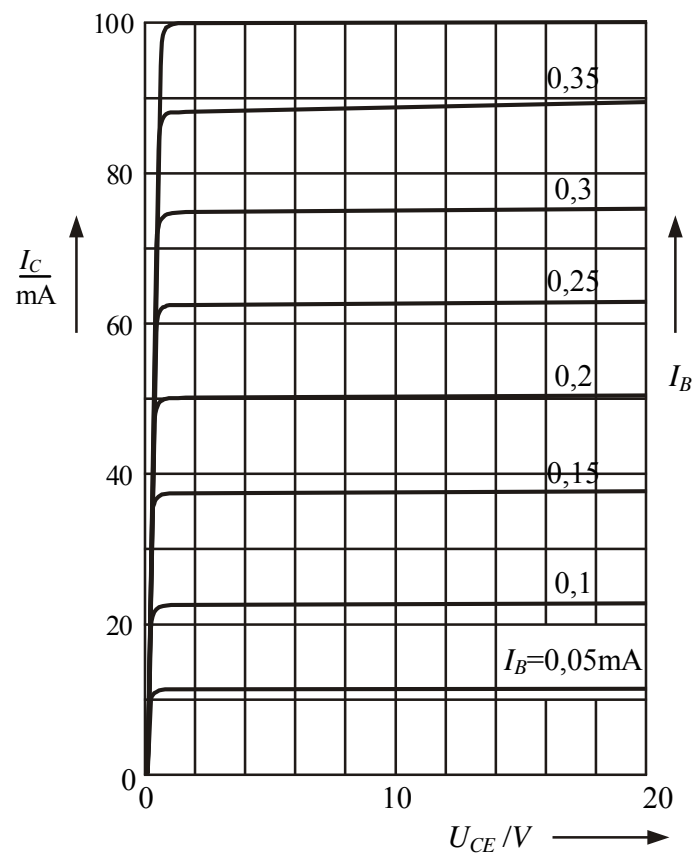


Bild 3.19: Ausgangskennlinienfeld eines bipolaren npn Transistors mit Ansteuerung über den Basisstrom I_B .

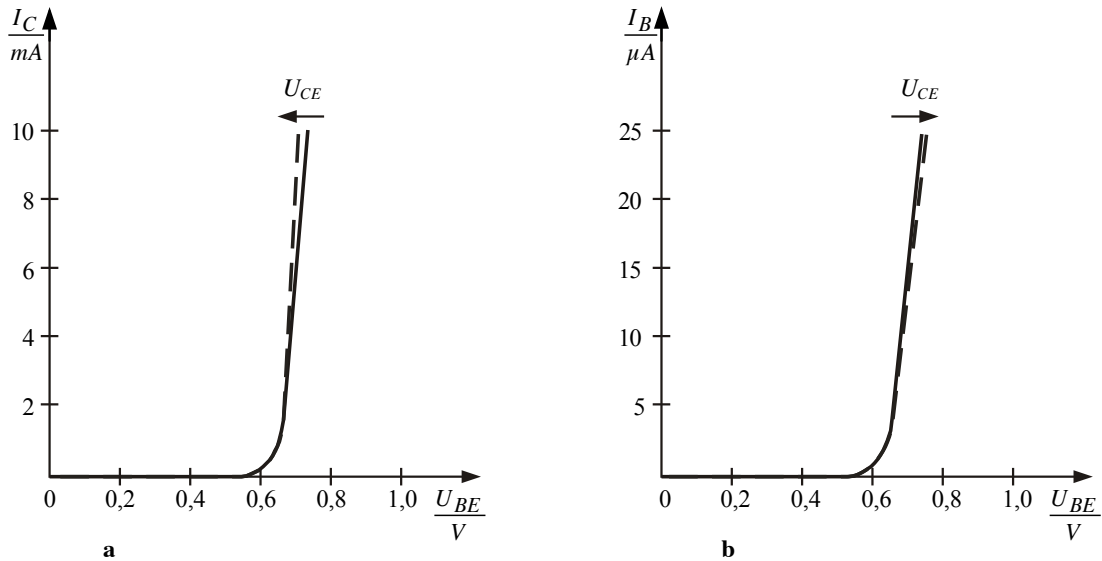


Bild 3.20: Eingangskennlinien von bipolaren Transistoren.

des Kennlinienfeldes. Die Kennlinien haben ungleiche Abstände, da der Kollektorstrom exponentiell mit (U_{BE}/U_T) ansteigt. Steuert man einen Transistor jedoch über den Basisstrom an, ist der Kollektorstrom I_C proportional dem Basisstrom und die Abstände der Kennlinien sind konstant. In Bild 3.19 ist ein solches Kennlinienfeld dargestellt.

Im Normalbetrieb ist der Kollektorstrom I_C im Wesentlichen nur von U_{BE} abhängig. Trägt man den Kollektorstrom in Abhängigkeit von der Basis-Emitterspannung für verschiedene Kollektor-Emitterspannungswerte auf, so erhält man das Eingangskennlinienfeld, das in Bild 3.20a gezeigt wird. Der Einfluss der Kollektor-Emitterspannung auf das Eingangskennlinienfeld ist sehr gering. Zur vollständigen Beschreibung wird noch das in Bild 3.20b gezeigte Eingangskennlinienfeld benötigt, bei dem der Basisstrom für verschiedene, zum Normalbetrieb gehörende Werte der Kollektor-Emitterspannung als Funktion der Basis-Emitterspannung aufgetragen ist. Auch hier ist der Einfluss der Kollektor-Emitterspannung sehr gering. Vergleicht man die beiden Kennlinien in Bild 3.20, d.h. die Übertragungskennlinie und die Eingangskennlinie, so sieht man, dass der Verlauf sehr ähnlich ist. Daraus ergibt sich, dass im Normalbetrieb der Kollektorstrom dem Basisstrom näherungsweise proportional folgt. Die Proportionalitätskonstante wird als Großsignal-Stromverstärkung B bezeichnet.

$$B = \frac{I_C}{I_B} \quad (3.9)$$

Die mathematische Beschreibung basiert darauf, dass das Verhalten des Bipolartransistors im Wesentlichen auf die Basis-Emitterdiode zurückgeführt werden kann. Der für eine Diode typische exponentielle Zusammenhang zwischen Strom und Spannung zeigt sich im Übertragungs- und im Eingangskennlinienfeld des Transistors als exponentielle

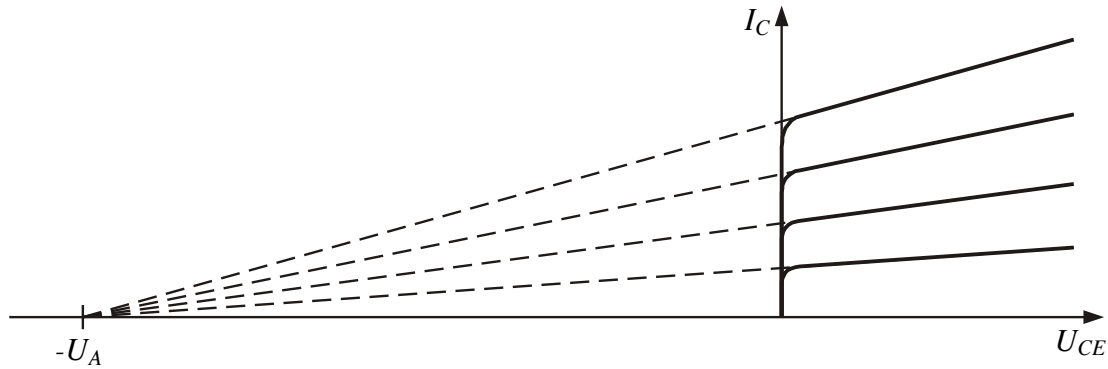


Bild 3.21: Early-Effekt.

Abhängigkeit der Ströme I_B und I_C von der Spannung U_{BE} . Damit ergibt sich folgender Zusammenhang:

$$I_C = B \cdot I_S \cdot e^{\frac{U_{BE}}{U_T}} \cdot \left(1 + \frac{U_{CE}}{U_A}\right) \quad (3.10)$$

Dabei gilt Gl. 3.9 und U_T als Temperaturspannung von 26 mV sowie $I_S \cong 10^{-16}$ bis 10^{-12} A als Sättigungssperrstrom des Transistors. U_A ist die Early-Spannung des Transistors.

3.2.2.2 Der Early-Effekt

Die Abhängigkeit der Kennlinien von U_{CE} wird durch den Early-Effekt verursacht. Dieser Effekt ist im rechten Term in Gleichung 3.10 empirisch beschrieben. Die Grundlage für diese Beschreibung ist die Beobachtung, dass sich die extrapolierten Kennlinien des Ausgangskennlinienfeldes näherungsweise in einem Punkt schneiden. Bild 3.21 zeigt diesen Effekt im Ausgangskennlinienfeld. Die Konstante U_A heißt Early-Spannung und beträgt bei npn Transistoren etwa 30 bis 150 V und bei pnp Transistoren 30 bis 75 V. Der Early-Effekt lässt sich darauf zurückführen, dass die Basis-Emitter- und die Basis-Kollektorspannung die effektive Dicke der Basiszone im Transistor beeinflussen und damit wiederum auf den Transportstrom I_C , also den Kollektorstrom, einwirken. In den meisten Fällen wird die Abhängigkeit der Stromverstärkung von der Basis-Emitter- bzw. der Kollektor-Emitterspannung vernachlässigt. Damit ist die Stromverstärkung B eine unabhängige Konstante. Um das Verhalten des Transistors genauer zu beschreiben, muss man den Early-Effekt berücksichtigen und damit gilt:

$$B = B_0 \cdot \left(1 + \frac{U_{CE}}{U_A}\right) \quad (3.11)$$

Typische Werte der Stromverstärkung, der sog. Großsignal-Stromverstärkung B , betragen 100 bis 500.

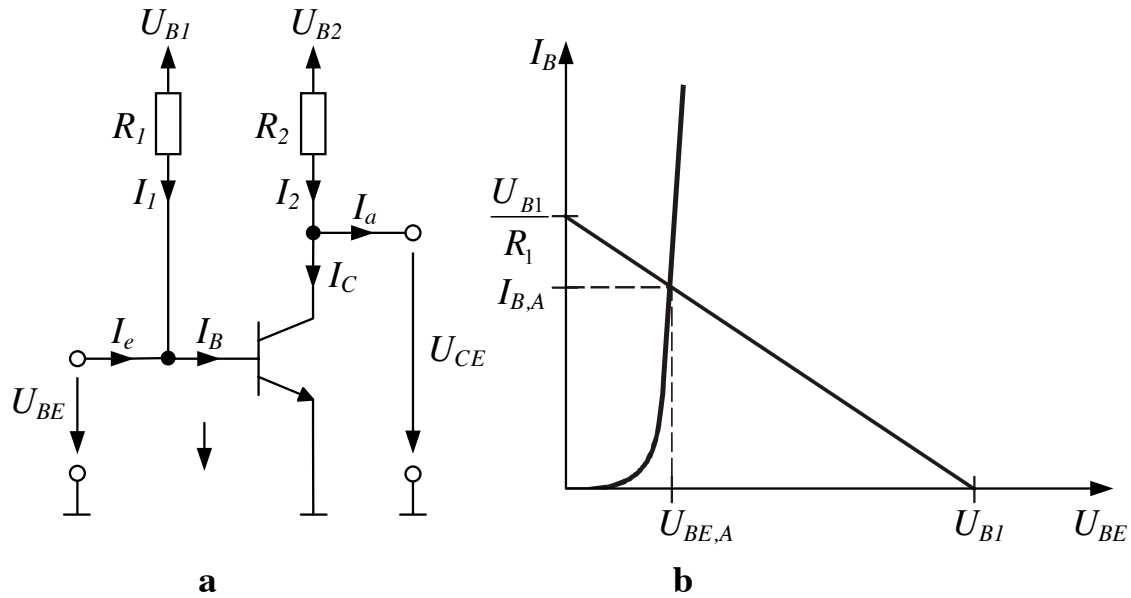


Bild 3.22: Einstellung des Arbeitspunktes.

3.2.3 Arbeitspunkt und Kleinsignal-Verhalten

3.2.3.1 Bestimmung des Arbeitspunktes

Ein wichtiges Anwendungsgebiet von Bipolartransistoren ist die Verstärkung von Signalen im Kleinsignalbetrieb, d.h. mit so kleinen Eingangssignalen, dass der Transistor lediglich in einem kleinen Bereich um den Arbeitspunkt A angesteuert wird. Der Arbeitspunkt A wird durch die Spannungen Kollektor-Emitter und Basis-Emitter, sowie dem Kollektor- und dem Basisstrom charakterisiert und durch die äußere Beschaltung des Transistors festgelegt. Jede Festlegung wird Arbeitspunkteinstellung genannt. In Bild 3.22 sind eine einfache Schaltung einer Transistor-Verstärkerstufe mit einem Transistor sowie das zugehörige Eingangskennlinienfeld schematisch dargestellt. Für $I_e = I_a = 0$ folgt aus dem Knotensatz folgendes Gleichungssystem:

$$\begin{aligned}
 I_C &= I_C(U_{BE}, U_{CE}) \\
 I_B &= I_B(U_{BE}, U_{CE}) \\
 I_B &= I_1 = \frac{U_{B1} - U_{BE}}{R_1} \\
 I_C &= I_2 = \frac{U_{B2} - U_{CE}}{R_2}
 \end{aligned} \tag{3.12}$$

Das Gleichungssystem 3.12 enthält 4 Unbekannte und ist damit lösbar. Neben der numerischen Lösung ist auch eine graphische Lösung möglich. Für die graphische Lösung zeichnet man in das Eingangskennlinienfeld und in das Ausgangskennlinienfeld die entsprechenden Lastgeraden, die aus den Gleichungen 3.12 berechnet werden können, ein und ermittelt die Schnittpunkte mit der entsprechenden Kennlinie der Transistors.

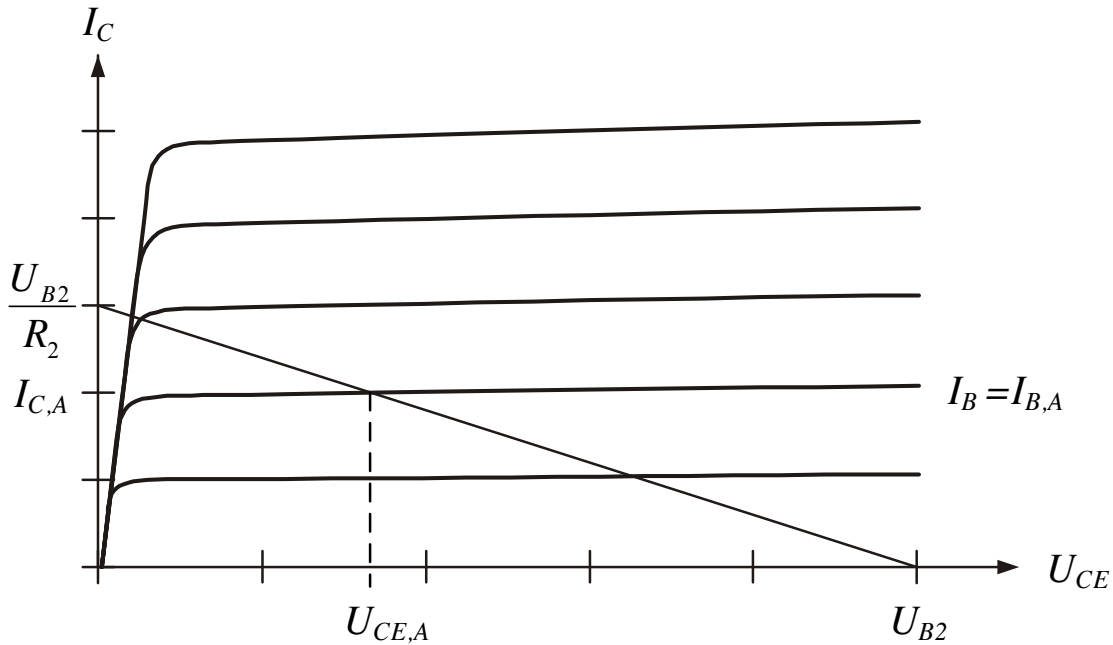


Bild 3.23: Arbeitspunkt im Ausgangskennlinienfeld.

Da das Eingangskennlinienfeld wegen der geringen Abhängigkeit von der Kollektor-Emitterspannung praktisch nur aus einer Kennlinie besteht, erhält man nach Bild 3.22b nur einen Schnittpunkt und kann damit direkt die Basis-Emitterspannung und den Basisstrom im Arbeitspunkt ablesen. Im Ausgangskennlinienfeld (Bild 3.23) kann man nun die Kollektor-Emitterspannung und den Kollektorstrom im Arbeitspunkt aus dem Schnittpunkt der Geraden mit der zu $I_{B,A}$ gehörigen Ausgangskennlinie bestimmen. Sowohl die analytische als auch die graphische Bestimmung des Arbeitspunktes sind so genannte analytische Verfahren, d.h. man kann damit bei bekannter Beschaltung einer Transistorstufe den Arbeitspunkt ermitteln. Zum Entwurf von neuen Schaltungen werden dagegen so genannte Syntheseverfahren benötigt, mit denen man dann die zu einem gewünschten Arbeitspunkt gehörende Beschaltung bestimmen kann. Diese Verfahren werden bei der Behandlung der verschiedenen Grundschaltungen von Bipolartransistoren behandelt.

3.2.3.2 Die Kleinsignal-Parameter

Für die Beschreibung und die Berechnung von Verstärkerschaltungen hat sich die Nutzung von Kleinsignal-Gleichungen und so genannten Kleinsignal-Parametern der Bipolartransistoren bewährt. Hierfür muss man die Kleinsignal-Größen, d.h. die Kleinsignal-Spannungen und Kleinsignal-Ströme um den Arbeitspunkt herum definieren:

$$\begin{aligned} u_{BE} &= U_{BE} - U_{BE,A} & i_B &= I_B - I_{B,A} \\ u_{CE} &= U_{CE} - U_{CE,A} & i_C &= I_C - I_{C,A} \end{aligned} \quad (3.13)$$

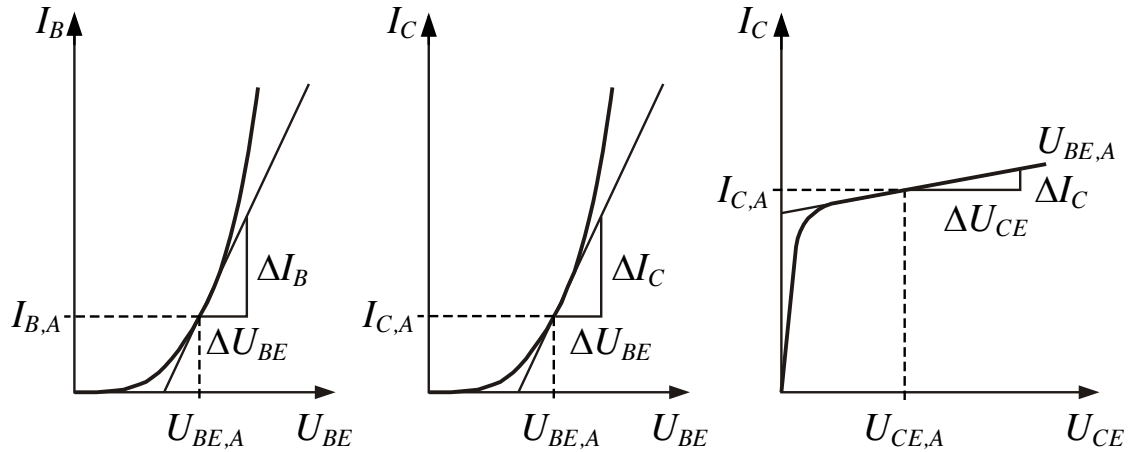


Bild 3.24: Linearisierung der Eingangs-, Übertragungs- und Ausgangskennlinien.

Um den Arbeitspunkt werden die Kennlinien der Transistoren linearisiert, d.h. die nichtlinearen Kennlinien werden im Arbeitspunkt durch ihre Tangenten ersetzt, was in Bild 3.24 am Beispiel der Übertragungskennlinie dargestellt ist. Um die Linearisierung zu verstehen, kann man als Beispiel den Basisstrom in Form einer Taylor-Reihe entwickeln. Es ergibt sich damit:

$$\begin{aligned}
 i_B &= I_B \cdot (U_{BE,A} + u_{BE}, U_{CE,A} + u_{CE}) - I_{B,A} = \\
 &= \left. \frac{\partial I_B}{\partial U_{BE}} \right|_A \cdot u_{BE} + \left. \frac{\partial I_B}{\partial U_{CE}} \right|_A \cdot u_{CE} + \dots
 \end{aligned} \tag{3.14}$$

Die partiellen Ableitungen im Arbeitspunkt werden dann Kleinsignalparameter genannt. Wenn man die entsprechenden partiellen Ableitungen wie folgt benennt, erhält man die so genannten Kleinsignal-Gleichungen des Bipolartransistors:

$$i_B = \frac{1}{r_{BE}} \cdot u_{BE} + S_r \cdot u_{CE} \tag{3.15}$$

$$i_C = S \cdot u_{BE} + \frac{1}{r_{CE}} \cdot u_{CE} \tag{3.16}$$

In diesen beiden Gleichungen sind folgende Kleinsignalparameter eingeführt worden: Die Steilheit S beschreibt die Änderung des Kollektorstromes I_C mit der Basis-Emitterspannung U_{BE} im Arbeitspunkt. Durch Differentiation der Gleichung 3.10 erhält man:

$$S = \frac{\partial I_C}{\partial U_{BE}} = \frac{I_{C,A}}{U_T} \tag{3.17}$$

Der Kleinsignalwiderstand r_{BE} beschreibt die Änderung der Basis-Emitterspannung

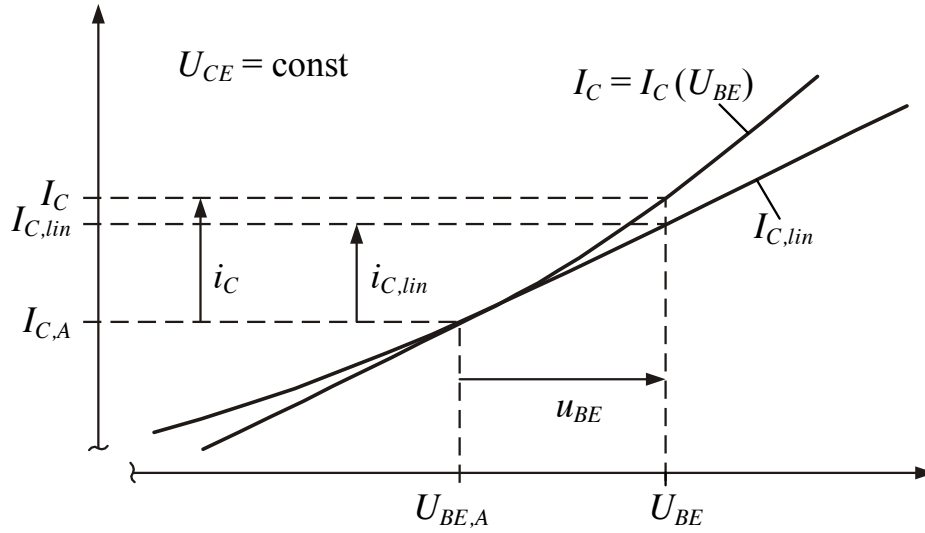


Bild 3.25: Bestimmung der Kleinsignalparameter aus den Kennlinien.

U_{BE} mit dem Basisstrom I_B im Arbeitspunkt. Mit Hilfe des folgenden Zusammenhanges

$$r_{BE} = \left. \frac{\partial U_{BE}}{\partial I_B} \right|_A = \left. \frac{\partial U_{BE}}{\partial I_C} \right|_A \cdot \left. \frac{\partial I_C}{\partial I_B} \right|_A \quad (3.18)$$

lässt sich der Kleinsignal-Eingangswiderstand wie folgt bestimmen:

$$r_{BE} = \frac{\beta}{S} \quad (3.19)$$

Der Kleinsignal-Ausgangswiderstand r_{CE} beschreibt die Änderung der Kollektor-Emitterspannung mit dem Kollektorstrom I_C im Arbeitspunkt. Er kann durch Differentiation der Gleichung 3.10 bestimmt werden.

$$r_{CE} = \left. \frac{\partial U_{CE}}{\partial I_C} \right|_A = \frac{U_A + U_{CE,A}}{I_{C,A}} \approx \frac{U_A}{I_{C,A}} \quad (3.20)$$

In der Praxis arbeitet man mit der in der Gleichung 3.20 angegebenen Näherung. Die Rückwärts-Steilheit S_r beschreibt die Änderung des Basisstromes I_B durch eine Änderung der Kollektor-Emitterspannung U_{CE} .

$$S_r = \frac{\partial I_B}{\partial U_{CE}} \approx 0 \quad (3.21)$$

Bild 3.25 zeigt die Möglichkeit auf, wie man die Kleinsignal-Parameter aus den Kennlinienfeldern bestimmt. Dieses Verfahren wird aber in der Praxis aufgrund der begrenzten Ablesegenauigkeit aus den Kennlinienfeldern sehr selten benutzt.

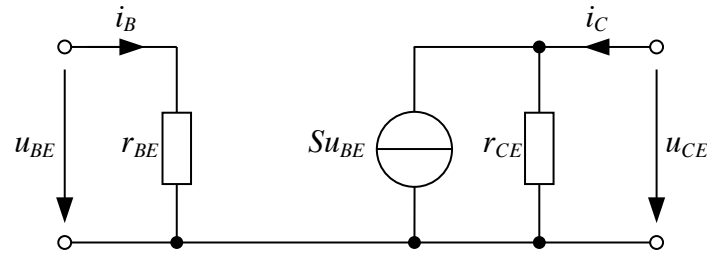


Bild 3.26: Kleinsignal-Ersatzschaltbild eines npn Transistors.

3.2.3.3 Wechselstrom-Kleinsignal-Ersatzschaltbild

Um das Kleinsignalverhalten von Transistorschaltungen im Bereich niedriger Frequenzen (0 bis 10 kHz) berechnen zu können, benötigt man das Kleinsignal-Ersatzschaltbild, was wir für diesen niedrigen Frequenzbereich auch als Wechselstrom-Ersatzschaltbild bezeichnen. Aus den Gleichungen 3.15 und 3.16 erhält man unter Vernachlässigung der Rückwirkung ($S_r = 0$) das in Bild 3.26 gezeichnete Kleinsignal-Ersatzschaltbild des Bipolartransistors. Die Gleichungen 3.15 und 3.16 können auch in Matrizen-Form geschrieben werden:

$$\begin{bmatrix} i_B \\ i_C \end{bmatrix} = \begin{bmatrix} \frac{1}{r_{BE}} & S_r \\ S & \frac{1}{r_{CE}} \end{bmatrix} \cdot \begin{bmatrix} u_{BE} \\ u_{CE} \end{bmatrix} \quad (3.22)$$

Diese Schreibweise entspricht der Leitwert-Darstellung eines Vierpols. Damit lässt sich die Leitwert-Matrix Y_e bestimmen:

$$\begin{bmatrix} i_B \\ i_C \end{bmatrix} = Y_e \cdot \begin{bmatrix} u_{BE} \\ u_{CE} \end{bmatrix} = \begin{bmatrix} Y_{11,e} & Y_{12,e} \\ Y_{21,e} & Y_{22,e} \end{bmatrix} \cdot \begin{bmatrix} u_{BE} \\ u_{CE} \end{bmatrix} \quad (3.23)$$

Die hier verwendeten Indizes „e“ deuten auf die Emitterschaltung hin. Die Matrizen-Darstellung kann natürlich für jede andere Grundschaltung ebenfalls aufgeschrieben werden. Weit verbreitet ist auch die so genannte Darstellung mit Strom-Aussteuerung (H-Matrix: Hybrid-Matrix, Bild 3.27). Die Darstellung mittels h-Parameter ist weit verbreitet, da sie auch in Datenblättern zur Charakterisierung der Kleinsignal-Parameter benutzt wird. Es gilt:

$$\begin{aligned} u_{BE} &= h_{11} \cdot i_B + h_{12} \cdot u_{CE} \\ i_C &= h_{21} \cdot i_B + h_{22} \cdot u_{CE} \end{aligned} \quad (3.24)$$

Aus der Analogie zu den Gleichungen 3.15 und 3.16 folgt die Umrechnung zwischen den Kleinsignal-Parametern, den h-Parametern und den Y-Parametern wie folgt:

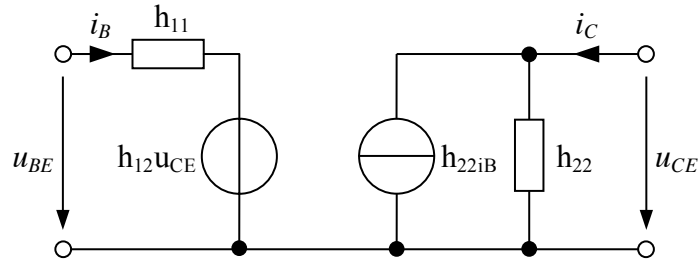


Bild 3.27: Ersatzschaltbild mit h-Parametern.

$$\begin{aligned}
 r_{BE} = h_{11,e} &= \frac{1}{Y_{11,e}} \\
 \beta = h_{21,e} &= \frac{Y_{21,e}}{Y_{11,e}} \\
 S &= \frac{h_{21,e}}{h_{11,e}} = Y_{21,e} \\
 S_r &= -\frac{h_{21,e}}{h_{11,e}} = Y_{12,e} \\
 r_{CE} &= \frac{h_{11,e}}{h_{11,e} \cdot h_{22,e} - h_{12,e} \cdot h_{21,e}} = \frac{1}{Y_{22,e}}
 \end{aligned} \tag{3.25}$$

Wie bereits erwähnt, gilt das Kleinsignal-Ersatzschaltbild nur für kleine Aussteuerungen um den Arbeitspunkt A . Als Maß für die Grenze der Anwendbarkeit oder der Aussteuerbarkeit eines Verstärkers kann man das Amplitudenverhältnis aus der Amplitude der 1. Oberwelle zur Amplitude der Grundwelle wählen. Wenn dieses Verhältnis kleiner als 1 % ist, kann man das Kleinsignal-Ersatzschaltbild anwenden.

3.2.3.4 Grenzdaten und zuverlässiger Arbeitsbereich

Bei einem Bipolartransistor werden verschiedene Grenzdaten angegeben, die nicht überschritten werden dürfen. Sie gliedern sich in Grenzspannungen, Grenzströme und die maximale Verlustleistung. Die Durchbruchspannungen (breakdown voltages) entsprechen den Durchbruchspannungen der einzelnen Diodenzweige im Transistor im Sperrgebiet. Sie werden durch den Zusatz „BR“ bezeichnet. Der Index „0“ gibt an, dass der 3. Anschluss offen (open) ist. Somit ergibt die Emitter-Basis-Durchbruchspannung $U_{(BR)EBO}$ die Durchbruchspannung der Emitterdiode im Sperrbetrieb an. Für fast alle Transistoren gilt $U_{(BR)EBO} \cong 5$ bis 7 V. Damit ist $U_{(BR)EBO}$ die kleinste Durchbruchspannung. Da $U_B < 0$ in der Praxis selten vorkommt, ist diese Durchbruchspannung von geringer Bedeutung. Die Kollektor-Basis-Durchbruchspannung $U_{(BR)CBO}$ bricht die Kollektordiode im Sperrbetrieb durch. Da im Normalbetrieb die Kollektordiode gesperrt ist, ist durch $U_{(BR)CBO}$ eine für die Praxis sehr wichtige Obergrenze gegeben. Sie beträgt bei

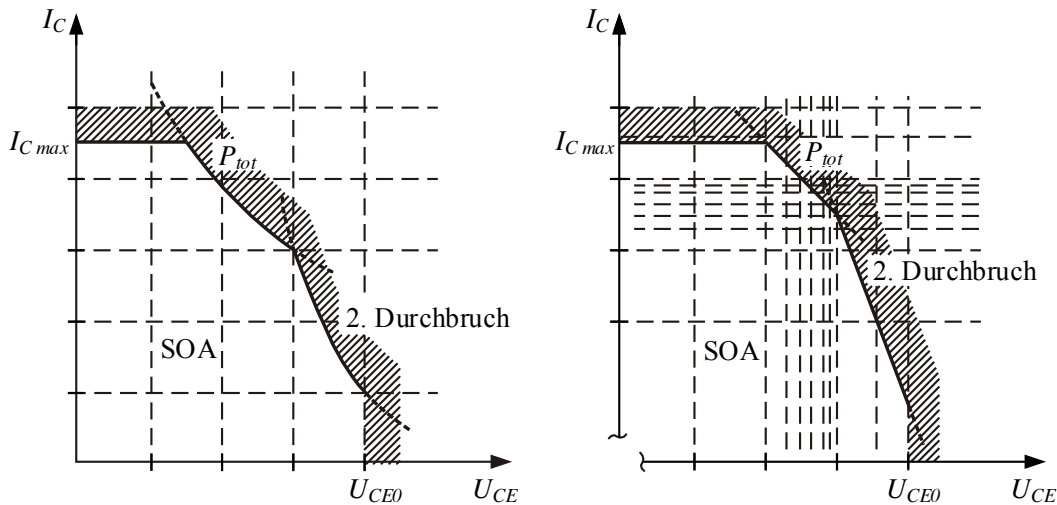


Bild 3.28: Zulässiger Betriebsbereich (safe operation area – SOA), (a) linear, (b) doppelt logarithmisch.

Niederspannungstransistoren etwa 20 bis 80 V, bei Hochspannungstransistoren erreicht sie Werte bis zu 1000 V und höher. Die maximale Kollektor-Emitterspannung ist für die Praxis sehr wichtig. Sie wird mit $U_{(BR)CE0}$ bezeichnet. Die maximalen Dauerströme werden in den Datenblättern mit dem Index „max“ bezeichnet. Die maximalen Spitzenwerte der Dauerströme für den gepulsten Betrieb sind etwa 2-mal größer als die maximalen Dauerströme. Eine besonders wichtige Grenzgröße ist die maximale Verlustleistung. Die Verlustleistung ist die im Transistor in Wärme umgesetzte Leistung:

$$P_V = U_{CE} \cdot I_C + U_{BE} \cdot I_B \approx U_{CE} \cdot I_C \quad (3.26)$$

Die maximale Verlustleistung führt zu einer Erwärmung der Sperrschicht der Kollektordiode. Diese Wärme kann nur über das Temperaturgefälle auf das Gehäuse an die Umgebung abgeführt werden. Die Temperatur der Sperrschicht in einem auf Silizium basierenden Transistor darf 175°C nicht überschreiten. In der Praxis wird mit einem Sicherheitswert von 150°C gerechnet. In den Datenblättern wird die maximale Verlustleistung je nach Transistor- und Gehäusotyp angegeben. Bei Kleinleistungstransistoren wird sie ohne Kühlkörper als P_{tot} für eine Umgebungstemperatur von 25°C angegeben. Für Leistungstransistoren wird die maximale Verlustleistung für eine bestimmte Temperatur des Gehäuses angegeben, gleichzeitig mit dem Wärmewiderstand des konkreten Transistortyps. Aus den Grenzdaten erhält man im Ausgangskennlinienfeld des Transistors den zulässigen Betriebsbereich (Safe Operating Area: SOA). Der SOA-Bereich wird durch den maximalen Kollektorstrom $I_{C,max}$, die Kollektor-/Emitter-Durchbruchspannung $U_{(BR)CE0}$ und die maximale Verlustleistung P_{tot} begrenzt. Bild 3.28 zeigt den SOA-Bereich. Die Kenntnis des SOA-Bereiches ist wichtig für die richtige Dimensionierung von Verstärkerstufen, insbesondere von Leistungsverstärkerstufen. Entsprechend Gl. 3.26 beträgt im statischen Betrieb die Verlustleistung

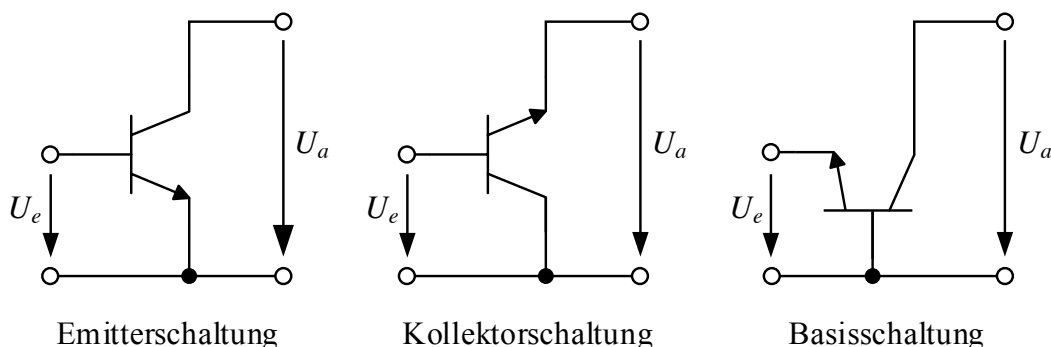


Bild 3.29: Die 3 Grundschaltungen eines Bipolartransistors.

$P_V = U_{CE} \cdot I_C$. Damit errechnet sich die Sperrschichttemperatur des Transistors wie folgt:

$$T_J = T_A + P_V \cdot R_{th,JA} \quad (3.27)$$

Dabei sind T_A die Umgebungstemperatur, P_V die Verlustleistung und $R_{th,JA}$ der Wärmewiderstand des Transistors zwischen der Sperrschicht und Umgebung. Der Wärmewiderstand ist definiert als $R_{th} = \frac{\Delta T}{P}$ und hat die Einheit $[\frac{K}{W}]$. Damit lässt sich die zulässige Verlustleistung im statischen Betrieb wie folgt berechnen:

$$P_{vmax,A} = \frac{T_{J,grenz} - T_{A,max}}{R_{th,JA}} \quad (3.28)$$

Mit Hilfe der Gleichung 3.28 ist somit je nach Umgebungstemperatur und je nach dem entsprechenden Wärmewiderstand der Kombination Transistor mit Kühlung die maximale Verlustleistung festgelegt.

3.2.4 Die Grundschaltungen

Bipolartransistoren können in 3 Grundschaltungen betrieben werden: der Emitterschaltung, der Kollektorschaltung und der Basisschaltung. Die Bezeichnung erfolgt jeweils entsprechend dem Anschluss des Transistors, der als gemeinsamer Bezugsknoten für den Eingang und den Ausgang der Schaltung dient. Bild 3.29 zeigt die verschiedenen Grundschaltungen. In der Praxis lässt sich kein einfacher Zusammenhang entsprechend der eben genannten Definition erfüllen, so dass man ein schwächeres Kriterium anwenden muss: Die Bezeichnung der Schaltung erfolgt entsprechend dem Anschluss des Transistors, der weder als Eingang noch als Ausgang der Schaltung dient. Auch in diesem Kapitel werden wir nur mit npn Transistoren arbeiten; sie können gegen pnp Transistoren ausgetauscht werden, wenn man die Polarität der Versorgungsspannungen und die Polarität der Elektrolytkondensatoren und Dioden entsprechend umtauscht.

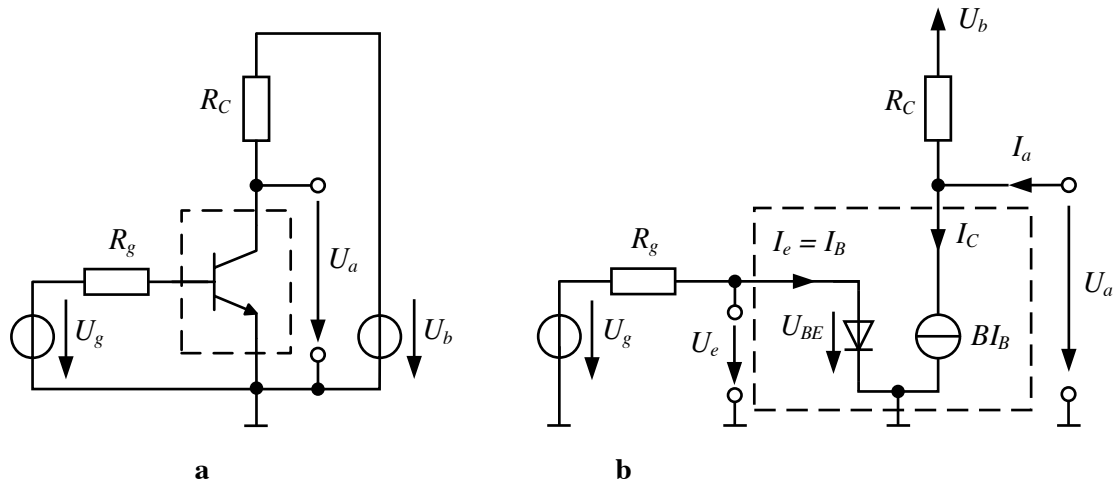


Bild 3.30: Emitterschaltung mit (a) Schalbild und (b) Großsignal-Ersatzschaltbild.

3.2.4.1 Die Emitterschaltung

Die einfachste Emitterschaltung ist in Bild 3.30a dargestellt. Sie besteht aus dem Transistor, dem Kollektorwiderstand R_C , der Signalspannungsquelle U_g , der Versorgungsspannung U_b mit dem Innenwiderstand der Eingangsspannungsquelle R_g . Die Ausgangsspannung wird mit U_a bezeichnet. Bild 3.30b zeigt das einfache Ersatzschaltbild für den Normalbetrieb in der Emitterschaltung. Um einfache Zahlenbeispiele rechnen zu können, soll gelten: $R_C = R_g = 1 \text{ k}\Omega$; $U_b = 5 \text{ V}$. Die Übertragungskennlinie ist in Bild 3.31 dargestellt. In Bild 3.31 sind 2 Bereiche gezeigt, der Bereich des Normalbetriebes und der Bereich des Sättigungsbetriebes. Der Transistor erreicht die Grenze zum Sättigungsbereich, wenn U_{CE} die Sättigungsspannung $U_{CE,sat}$ erreicht.

Normalbetrieb

Das Ersatzschaltbild für den Normalbetrieb ist in Bild 3.30b dargestellt. Wenn man den Early-Effekt vernachlässigt, gilt:

$$I_C = B \cdot I_B = B \cdot I_S \cdot e^{\frac{U_{BE}}{U_T}} \quad (3.29)$$

Damit erhält man für die Spannungen:

$$U_a = U_{CE} = U_b + (I_a - I_C) \cdot R_C = (U_b - I_C \cdot R_C) \Big|_{I_a=0} \quad (3.30)$$

$$U_e = U_{BE} = U_g - I_B \cdot R_g = U_g - \frac{I_C \cdot R_g}{B} \approx U_g \quad (3.31)$$

Als Arbeitspunkt wird ein Punkt in der Mitte des abfallenden Bereiches der Übertragungskennlinie gewählt. Damit wird die Aussteuerbarkeit maximal.

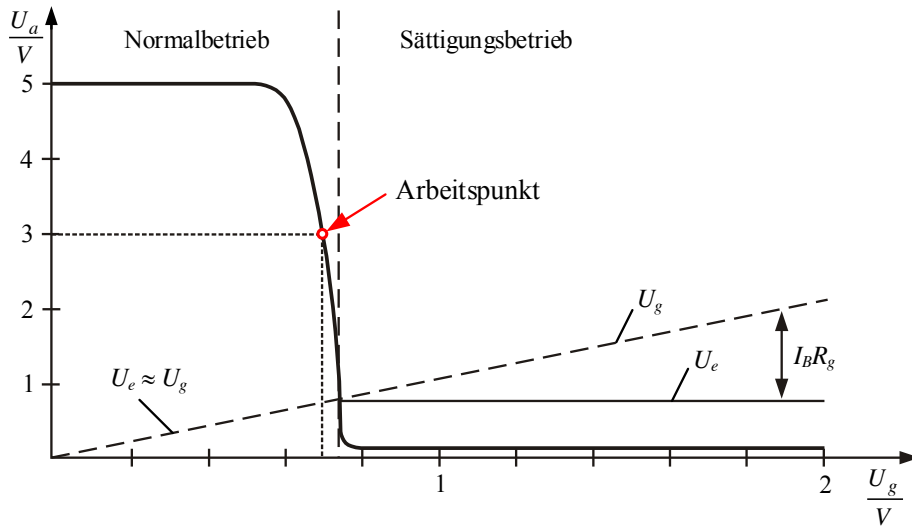


Bild 3.31: Übertragungskennlinie der Emitterschaltung.

Kleinsignalverhalten der Emitterschaltung

Zur Verdeutlichung des Zusammenhanges zwischen den nichtlinearen Kennlinien und dem Kleinsignalersatzschaltbild wird das Kleinsignalverhalten einmal aus den Kennlinien abgeleitet und anschließend aus dem Kleinsignalersatzschaltbild berechnet. Die Kleinsignaldarstellung für den einzelnen Transistor wurde in den Gleichungen 3.13 und 3.14 eingeführt. Für die Beschreibung der Schaltungen benötigen wir noch folgende zusätzliche Parameter. Die Kleinsignal-Spannungsverstärkung A entspricht der Steigung der Übertragungskennlinie, siehe Bild 3.32. Durch Differentiation von Gl. 3.30 erhält man:

$$A = \left. \frac{\partial u_a}{\partial u_e} \right|_A = \left. \frac{\partial i_C}{\partial u_{BE}} \right|_A \cdot R_C = -\frac{I_{C,A} \cdot R_C}{U_T} = -S \cdot R_C \quad (3.32)$$

Der Kleinsignal-Eingangswiderstand r_e ergibt sich aus der Eingangskennlinie:

$$r_e = \left. \frac{\partial u_e}{\partial i_e} \right|_A = \left. \frac{\partial u_{BE}}{\partial i_B} \right|_A = r_{BE} \quad (3.33)$$

Der Kleinsignal-Ausgangswiderstand r_a und der differentielle Kollektor-Emitterwiderstand r_{CE} sind:

$$r_a = \left. \frac{\partial u_a}{\partial i_a} \right|_A = R_C \quad r_{CE} = \left. \frac{\partial u_{CE}}{\partial i_C} \right|_A = r_{CE} \quad (3.34)$$

Durch Einsetzen des Ersatzschaltbildes des Transistors nach Bild 3.26 erhält man das Kleinsignal-Ersatzschaltbild der Emitterschaltung entsprechend Bild 3.33. Aus dem Kleinsignal-Ersatzschaltbild lassen sich ohne Kennlinien die Kleinsignalparameter berechnen. Man erhält die gleichen Parameter wie aus den Kennlinienfeldern, wenn man

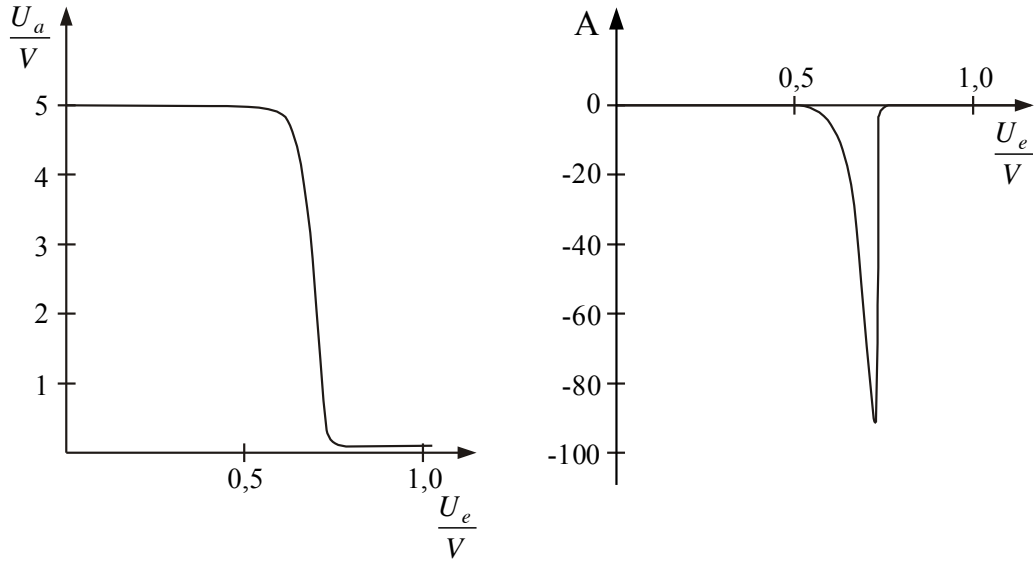


Bild 3.32: Übertragungskennlinie (links) und Verstärkung (rechts) der Emitterschaltung.

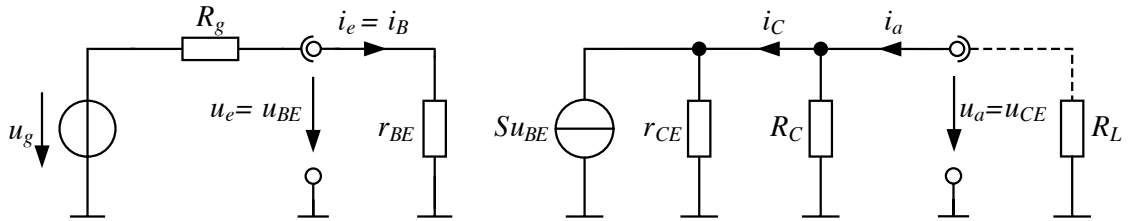


Bild 3.33: Kleinsignal-Ersatzschaltbild der Emitterschaltung.

den Early-Effekt vernachlässigt. Mit der Verstärkung A , den Eingangs- und Ausgangswiderständen r_e und r_a lässt sich eine Emitterschaltung vollständig beschreiben. Wenn man einen Lastwiderstand R_L anschließt, darf der Arbeitspunkt nicht verschoben werden, d.h. es darf nur ein vernachlässigbarer Gleichstrom über R_L fließen. Aus Bild 3.33 ergeben sich somit folgende Parameter der Emitterschaltung:

$$A = \frac{u_a}{u_e} \bigg|_{i_a=0} = -S \cdot (R_C \parallel r_{CE}) \stackrel{r_{CE} \gg R_C}{\approx} -S \cdot R_C \quad (3.35)$$

$$r_e = \frac{u_e}{i_e} = r_{BE} \quad (3.36)$$

$$r_a = \frac{u_a}{i_a} = R_C \parallel r_{CE} \stackrel{r_{CE} \gg R_C}{\approx} R_C \quad (3.37)$$

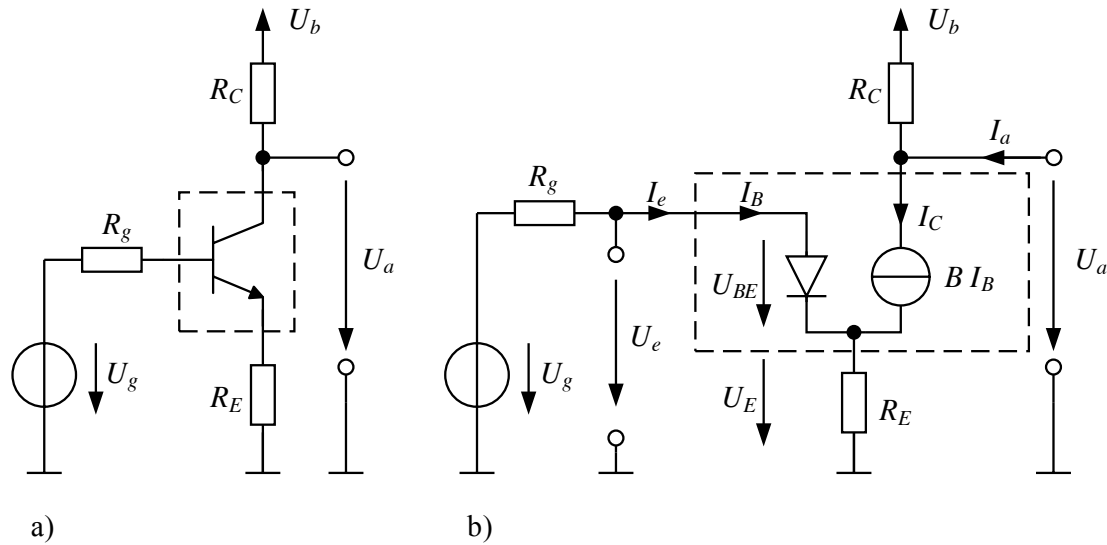


Bild 3.34: a) Emitterschaltung mit Stromgegenkopplung und b) Großsignal-Ersatzschaltbild.

Die Temperaturabhängigkeit des Arbeitspunktes ergibt sich aus der Tatsache, dass die Basis-Emitterspannung U_{BE} bei konstantem Kollektorstrom I_C mit 2 mV/K abnimmt. Man muss also die Eingangsspannung um 2 mV/K verringern, damit der Kollektorstrom konstant bleibt. Hält man hingegen die Eingangsspannung konstant, wirkt sich eine Temperaturerhöhung wie eine Zunahme der Eingangsspannung aus. Somit kann man die Temperaturdrift wie folgt berechnen:

$$\left. \frac{\partial U_a}{\partial T} \right|_A = \left. \frac{\partial U_a}{\partial U_e} \right|_A \cdot \frac{\partial U_e}{\partial T} \approx A \cdot 2 \frac{\text{mV}}{\text{K}} \quad (3.38)$$

Bei einer Verstärkung von z.B. $A = -100$ ergibt sich eine Temperaturdrift von -200 mV/K . Mit der Temperaturdrift verschieben sich A , r_e , r_a . Um einen stabilen Betrieb der Emitterschaltung zu ermöglichen, sind entsprechende Maßnahmen erforderlich. Eine besonders effektive Methode stellt die Gegenkopplung dar.

Emitterschaltung mit Stromgegenkopplung

Die Temperaturdrift und die Nichtlinearitäten der Emitterschaltung kann durch eine Stromgegenkopplung verringert werden. Entsprechend Bild 3.34 wird ein Emitterwiderstand R_E eingebaut. Rechts daneben ist das Großsignal-Ersatzschaltbild dargestellt. Man erhält folgende Spannungen:

$$U_a = U_b + (I_a - I_C) \cdot R_C \stackrel{I_a=0}{=} U_b - I_C \cdot R_C \quad (3.39)$$

$$U_e = U_{BE} + U_E = U_{BE} + (I_C + I_B) \cdot R_E \approx U_{BE} + I_C \cdot R_E \quad (3.40)$$

$$U_e = U_g - I_B \cdot R_g \approx U_g \quad (3.41)$$

Für eine Stromverstärkung $B \gg 1$, kann der Basisstrom in Gleichung 3.40 gegenüber dem Kollektorstrom vernachlässigt werden. Weiterhin soll der Spannungsabfall an R_g vernachlässigbar sein. Die Stromgegenkopplung zeigt sich dadurch, dass die Basis-Emitterspannung um den Betrag $\Delta U_{BE} = \Delta I_C \cdot R_E$ gegenüber der nicht gegengekoppelten Emitterschaltung verringert wird.

Kleinsignalverhalten der Emitterschaltung mit Stromgegenkopplung

Die Spannungsverstärkung A berechnet sich aus dem Kleinsignal-Ersatzschaltbild Bild 3.35 mit Hilfe des Knotensatzes:

$$\frac{u_e - u_E}{r_{BE}} + S \cdot u_{BE} + \frac{u_a - u_E}{r_{CE}} = \frac{u_E}{R_E} \quad (3.42)$$

$$S \cdot u_{BE} + \frac{u_a - u_E}{r_{CE}} + \frac{u_a}{R_C} = i_a \quad (3.43)$$

Mit $u_{BE} = u_e - u_E$ erhält man:

$$A = \left. \frac{u_a}{u_e} \right|_{i_a=0} = - \frac{S \cdot R_C \cdot (1 - \frac{R_E}{\beta \cdot r_{CE}})}{1 + R_E \cdot (S \cdot (1 + \frac{1}{\beta} + \frac{R_C}{\beta \cdot r_{CE}}) + \frac{1}{r_{CE}}) + \frac{R_C}{r_{CE}}} \quad (3.44)$$

$$\stackrel{r_{CE} \gg R_C, R_E}{\approx} \stackrel{\beta \gg 1}{\approx} - \frac{S \cdot R_C}{1 + S \cdot R_E} \stackrel{s R_E \gg 1}{\approx} - \frac{R_C}{R_E}$$

Damit hängt für $S \cdot R_E \gg 1$ die Verstärkung nur noch von R_C und R_E ab. Falls ein Lastwiderstand R_L angeschlossen ist, kann man die zugehörige Verstärkung einfach mittels Substitution den Widerstandes R_C durch die Parallelschaltung von R_C und R_L ersetzt. Für den Eingangswiderstand erhält man für $i_a = 0$

$$r_e = \left. \frac{u_e}{i_e} \right|_{i_a=0} = r_{BE} + \frac{(1 + \beta) \cdot r_{CE} + R_C}{r_{CE} + R_E + R_C} \stackrel{r_{CE} \gg R_C, R_E}{\approx} \stackrel{\beta \gg 1}{\approx} r_{BE} + \beta \cdot R_E \quad (3.45)$$

Die Abhängigkeit vom Lastwiderstand ist sehr gering und kann durch die Parallelschaltung von R_C und R_L berücksichtigt werden. Der Ausgangswiderstand wird auch durch

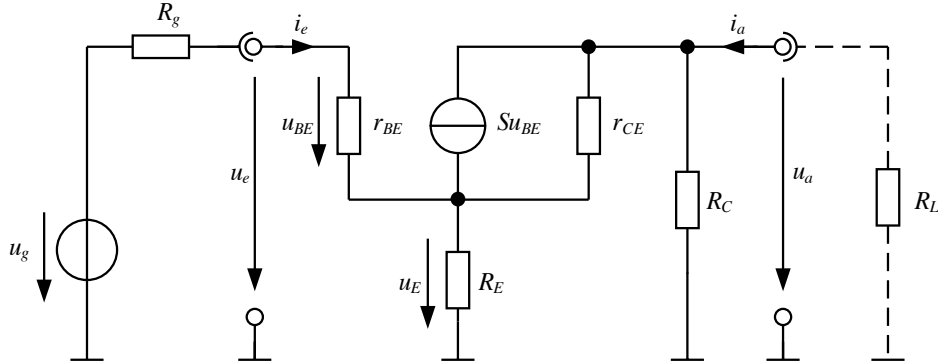


Bild 3.35: Kleinsignal-Ersatzschaltbild der Emitterschaltung mit Stromgegenkopplung.

den Innenwiderstand der Quelle R_g beeinflusst. Deshalb betrachten wir hier nur Grenzfälle. Der Kurzschlussausgangswiderstand gilt für einen Kurzschluss am Eingang, d.h. $u_e = 0$ bzw. $R_g = 0$ und ergibt sich entsprechend nach folgender Gleichung.

$$r_{a,K} = \left. \frac{u_a}{i_a} \right|_{u_e=0} = R_C \parallel r_{CE} \cdot \left(1 + \frac{\beta + \frac{r_{BE}}{r_{CE}}}{1 + \frac{r_{BE}}{R_E}} \right) \quad (3.46)$$

$$\begin{aligned} & r_{CE} \gg r_{BE} \\ & \beta \gg 1 \\ & \approx R_C \parallel r_{CE} \cdot \frac{\beta \cdot R_E + r_{BE}}{R_E + R_{BE}} \stackrel{r_{CE} \gg R_C}{\approx} R_C \end{aligned}$$

Der Leerlaufausgangswiderstand ergibt sich mit $i_e = 0$ bzw. $R_g \rightarrow \infty$.

$$r_{a,L} = \left. \frac{u_a}{i_a} \right|_{i_e=0} = R_C \parallel (R_E + r_{CE}) \stackrel{r_{CE} \gg R_C}{\approx} R_C \quad (3.47)$$

Auch hier ist der Einfluss von R_g sehr gering.

Zusammenfassend erhält man mit $r_{CE} \gg R_C, R_E$; $\beta \gg 1$ und $R_L \rightarrow \infty$ folgende Gleichungen der Parameter der Emitterschaltung mit Stromgegenkopplung:

$$A = \left. \frac{u_a}{u_e} \right|_{i_a=0} \approx -\frac{S \cdot R_C}{1 + S \cdot R_E} \stackrel{S R_E \gg 1}{\approx} -\frac{R_C}{R_E} \quad (3.48)$$

$$r_e = \frac{u_e}{i_e} \approx r_{BE} + \beta \cdot R_E = r_{BE} \cdot (1 + S \cdot R_E) \quad (3.49)$$

$$r_a = \frac{u_a}{i_a} \approx R_C \quad (3.50)$$

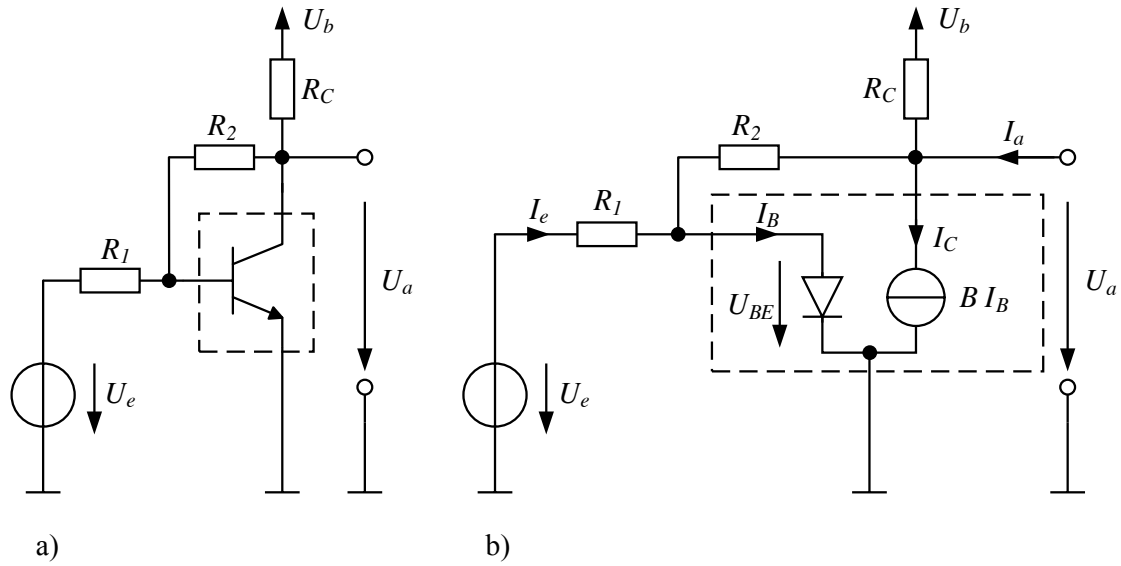


Bild 3.36: Emitterschaltung mit Spannungsgegenkopplung im Großsignal-Ersatzschaltbild.

Der Vergleich zu der Emitterschaltung ohne Stromgegenkopplung zeigt, dass die Verstärkung um den Faktor $(1 + S \cdot R_C)$ kleiner wird und der Eingangswiderstand um den gleichen Faktor zunimmt. Die Temperaturdrift verringert sich ebenfalls um den gleichen Faktor. Werte von einigen mV/K sind üblich.

Emitterschaltung mit Spannungsgegenkopplung

Bild 3.36 zeigt eine Emitterschaltung mit Spannungsgegenkopplung. Über die Widerstände R_2 und R_1 wird ein Teil der Ausgangsspannung auf die Basis des Transistors zurückgekoppelt. Aus Bild 3.36 ergibt sich mit dem Knotensatz

$$\frac{U_e - U_{BE}}{R_1} + \frac{U_a - U_{BE}}{R_2} = I_B = \frac{I_C}{B} \quad (3.51)$$

$$\frac{U_b - U_a}{R_C} + I_a = \frac{U_a - U_{BE}}{R_2} + I_C \quad (3.52)$$

Für $I_a = 0$ ergibt sich

$$U_a = \frac{U_b \cdot R_2 - I_C \cdot R_C \cdot R_2 + U_{BE} \cdot R_C}{R_2 + R_C} \quad (3.53)$$

$$U_e = \frac{I_C \cdot R_1}{B} + U_{BE} \cdot \left(1 + \frac{R_1}{R_2}\right) - U_a \cdot \frac{R_1}{R_2} \quad (3.54)$$

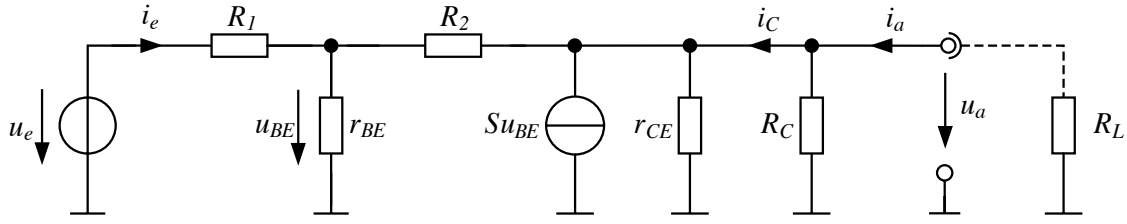


Bild 3.37: Kleinsignal-Ersatzschaltbild der Emitterschaltung mit Spannungsgegenkopplung.

Löst man Gleichung 3.53 nach I_C auf und setzt in Gleichung 3.54 ein folgt für $B \gg 1$ und $B \cdot R_C \gg R_2$:

$$U_a \approx \frac{U_b \cdot R_2}{B \cdot R_C} + \left(1 + \frac{R_2}{R_1}\right) \cdot U_{BE} - \frac{R_2}{R_1} \cdot U_e \quad (3.55)$$

Die Spannungsgegenkopplung bewirkt, dass die Verstärkung im Wesentlichen nur noch von R_1 und R_2 abhängt.

Kleinsignalverhalten der Emitterschaltung mit Spannungsgegenkopplung

Die Spannungsverstärkung A lässt sich für den Geltungsbereich der Gleichung 3.55 aus dem Kleinsignalersatzschaltbild nach Bild 3.37 herleiten. Aus dem Knotensatz folgt:

$$\frac{u_e - u_{BE}}{R_1} + \frac{u_a - u_{BE}}{R_2} = \frac{u_{BE}}{r_{BE}} \quad (3.56)$$

$$S \cdot u_{BE} + \frac{u_a - u_{BE}}{R_2} + \frac{u_a}{r_{CE}} + \frac{u_a}{R_C} = i_a \quad (3.57)$$

Mit $R'_C = R_C \parallel r_{CE}$ gilt:

$$\begin{aligned} A = \frac{u_a}{u_e} \bigg|_{i_a=0} &= - \frac{S \cdot R_2 + 1}{1 + R_1 \cdot \left(S \cdot \left(1 + \frac{1}{\beta}\right) + \frac{1}{R'_C} \right) + \frac{R_2}{R'_C} \cdot \left(1 + \frac{R_1}{r_{BE}}\right)} \\ &\stackrel{r_{CE} \gg R_C}{\approx} - \frac{S \cdot R_2 + 1}{1 + S \cdot R_1 + \frac{R_1}{R_C} + \frac{R_2}{R_C} \cdot \left(1 + \frac{R_1}{r_{BE}}\right)} \\ &\stackrel{\beta \gg 1}{\approx} - \frac{S \cdot R_2 + 1}{1 + S \cdot R_1 + \frac{R_1}{R_C} + \frac{R_2}{R_C} \cdot \left(1 + \frac{R_1}{r_{BE}}\right)} \\ &\stackrel{r_{BE} \gg R_1}{\approx} - \frac{R_2}{R_1 + \frac{R_1 + R_2}{S \cdot R_C}} \stackrel{SR_C \gg 1 + R_1/R_2}{\approx} - \frac{R_2}{R_1} \end{aligned} \quad (3.58)$$

Für den Leerlaufeingangswiderstand gilt mit $R'_C = R_C \parallel r_{CE}$:

$$\begin{aligned}
 r_{e,L} = \frac{u_e}{i_e} \Big|_{i_a=0} &= R_1 + \frac{r_{BE} \cdot (R'_C + R_2)}{r_{BE} + (1 + \beta) \cdot R'_C + R_2} \\
 &\stackrel{r_{CE} \gg R_C}{\approx} R_1 + \frac{r_{BE} \cdot (R_C + R_2)}{r_{BE} + \beta \cdot R_C + R_2} \\
 &\stackrel{\beta \gg 1}{\approx} R_1 + \frac{1}{S} \cdot \left(1 + \frac{R_2}{R_C} \right) \\
 &\stackrel{\beta R_C \gg r_{BE}, R_2}{\approx} R_1 + \frac{1}{S} \stackrel{SR_C \gg R_2/R_1}{\approx} R_1 + \frac{1}{S} \stackrel{SR_1 \gg 1}{\approx} R_1
 \end{aligned} \tag{3.59}$$

Den Kurzschlusseingangswiderstand erhält man durch $R_L = R_C = 0$

$$r_{e,K} = \frac{u_e}{i_e} \Big|_{u_a=0} = R_1 + r_{BE} \parallel R_2 \tag{3.60}$$

Der Kurzschlussausgangswiderstand ergibt sich mit $R'_C = R_C \parallel r_{CE}$ zu

$$\begin{aligned}
 r_{a,K} = \frac{u_a}{i_a} \Big|_{u_e=0} &= R'_C \parallel \frac{r_{BE} \cdot (R_1 + R_2) + R_1 \cdot R_2}{r_{BE} + R_1 \cdot (1 + \beta)} \\
 &\stackrel{r_{CE} \gg R_C}{\approx} R_C \parallel \frac{r_{BE} \cdot (R_1 + R_2) + R_1 \cdot R_2}{r_{BE} + \beta \cdot R_1} \\
 &\stackrel{\beta \gg 1}{\approx} R_C \parallel \left(\frac{1}{S} \cdot \left(1 + \frac{R_2}{R_1} \right) + \frac{R_2}{\beta} \right)
 \end{aligned} \tag{3.61}$$

Mit $R_L \rightarrow \infty$ ergibt sich der Leerlaufausgangswiderstand:

$$r_{a,L} = \frac{u_a}{i_a} \Big|_{i_e=0} = R'_C \parallel \frac{r_{BE} + R_2}{1 + \beta} \stackrel{r_{CE} \gg R_C}{\approx} R_C \parallel \left(\frac{1}{S} + \frac{R_2}{\beta} \right) \tag{3.62}$$

Zusammenfassend ergeben sich folgende Parameter für die Emitterschaltung mit Spannungsgegenkopplung:

$$A = \frac{u_a}{u_e} \Big|_{i_a=0} = - \frac{R_2}{R_1 + \frac{R_1 + R_2}{S \cdot R_C}} \stackrel{SR_C \gg 1 + R_2/R_1}{\approx} - \frac{R_2}{R_1} \tag{3.63}$$

$$r_e = \frac{u_e}{i_e} = R_1 \tag{3.64}$$

$$r_a = \frac{u_a}{i_a} = R_C \parallel \left(\frac{1}{S} \cdot \left(1 + \frac{R_2}{R_1} \right) + \frac{R_2}{\beta} \right) \quad (3.65)$$

Arbeitspunkteinstellung bei Wechselspannungskopplung

Bei Wechselspannungskopplung werden der Eingang und der Ausgang entsprechend mit der Signalquelle bzw. Last über einen Koppelkondensator verbunden (siehe Bild 3.38). Damit kann man die Gleichspannungen bzw. Gleichströme zur Arbeitspunkteinstellung unabhängig von den angeschlossenen Signalquellen und der Last einstellen. Weiterhin kann man damit bei mehrstufigen Verstärkern die Arbeitspunkte für jede Stufe unabhängig voneinander einstellen. Es ist dabei zu beachten, dass jeder Koppelkondensator einen Hochpass bildet. Die Dimensionierung der Kondensatoren muss so erfolgen, dass die kleinste zu übertragende Frequenz noch übertragen wird. Wir haben gezeigt, dass die Temperaturdrift proportional zur Verstärkung der Emitterschaltung ist. Damit kann man die Temperaturdrift durch Reduzierung der Verstärkung verkleinern. Da die Temperaturdrift sehr langsam verläuft, kann man die Gleichspannungsverstärkung reduzieren und die Wechselspannungsverstärkung unverändert lassen. Dazu wird nach Bild 3.39 ein Kondensator über den Emittewiderstand R_E geschaltet, der mit zunehmender Frequenz die Gegenkopplung aufhebt. Aus Gl. 3.48 folgt, dass man R_E möglichst groß machen sollte, damit die Gleichspannungsverstärkung gering ist und damit eine geringe Temperaturdrift zu erhalten. In der Praxis werden Werte von $R_C/R_E \approx 1 \dots 10$ gewählt. Der Frequenzgang der Gegenkopplung kann dann durch Einsetzen der Parallelschaltung $R_E \parallel (\frac{1}{\omega \cdot C})$ berechnet werden. Die Knickfrequenz dieses Hochpasses beträgt

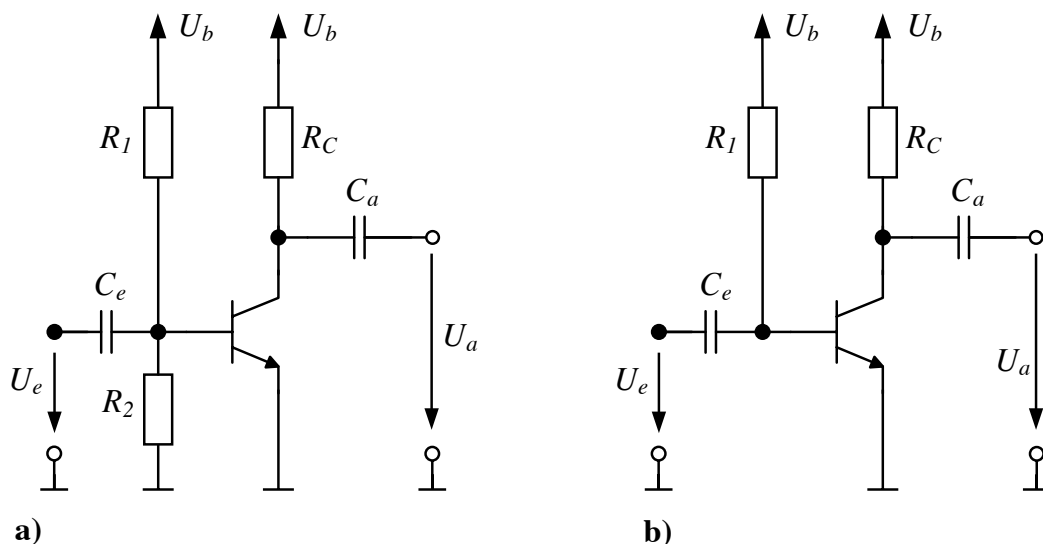


Bild 3.38: Arbeitspunkteinstellung bei Wechselspannungskopplung.

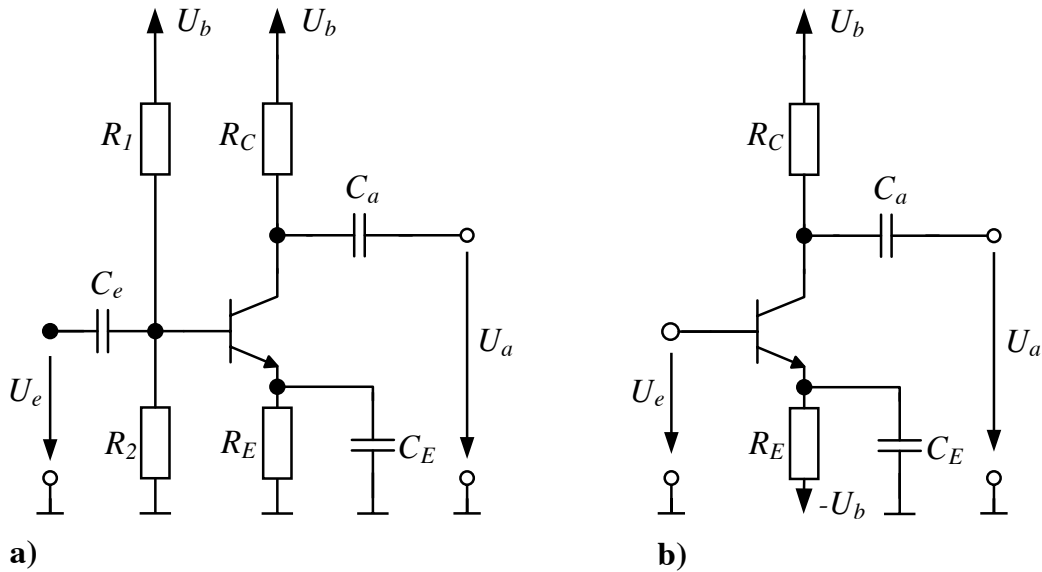


Bild 3.39: Arbeitspunkteinstellung bei Wechsellspannungskopplung mit Spannungseinstellung (a) und direkter Kopplung (b).

$f_E = \frac{1}{2\pi R_E C_E}$. Damit ist für Frequenzen $f < f_E$ die Gegenkopplung voll wirksam und für $f \gg f_E$ unwirksam.

3.2.4.2 Die Kollektorschaltung

Die Kollektorschaltung ist in Bild 3.40 dargestellt. Sie besteht aus einem Transistor, einem Emitterwiderstand R_E , der Versorgungsspannung U_b und der Signalspannungsquelle U_g mit dem Innenwiderstand R_g . Zur einfachen Beschreibung von Beispielen benutzen wir $U_b = 5 \text{ V}$, $R_E = R_g = 1 \text{ k}\Omega$. Das Übertragungsverhalten berechnen wir mit dem vereinfachten Ersatzschaltbild nach Bild 3.40. Der Kollektorstrom ist gegeben mit

$$I_C = B \cdot I_B = B \cdot I_S \cdot e^{\frac{U_{BE}}{U_T}} \quad (3.66)$$

Aus dem Knotensatz folgt

$$U_a = (I_C + I_B + I_a) \cdot R_E \approx (I_C + I_a) \cdot R_E \stackrel{I_a=0}{=} I_C \cdot R_E \quad (3.67)$$

$$U_e = U_a + U_{BE} \quad (3.68)$$

$$U_e = U_g - I_B \cdot R_g = U_g - \frac{I_C \cdot R_g}{B} \approx U_g \quad (3.69)$$

Dabei wurde angenommen, dass die Spannungsverstärkung B groß und R_g hinreichend klein ist, damit der Spannungsabfall an R_g vernachlässigt werden kann. In Gl. 3.67 wird der Basisstrom vernachlässigt. Aus Gleichung 3.68 erhält man unter der Annahme von

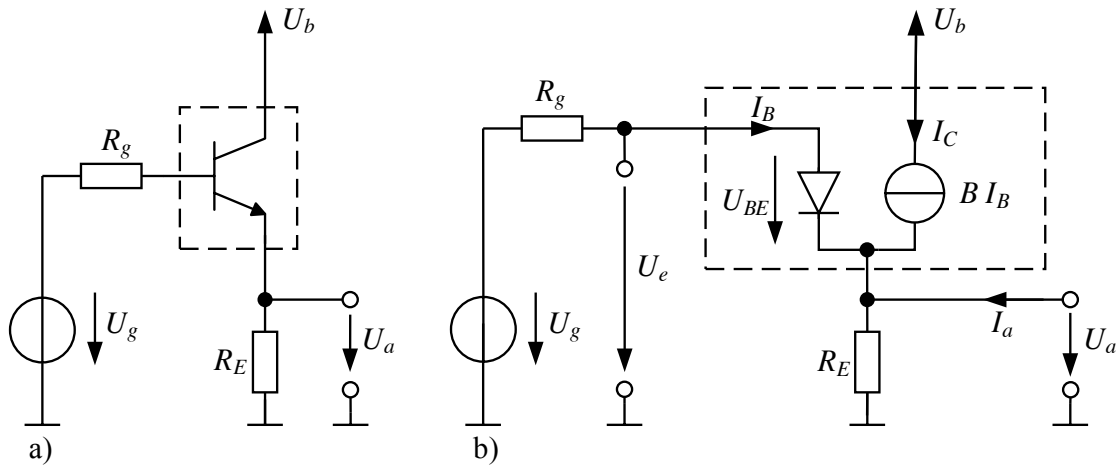


Bild 3.40: Kollektorschaltung mit Schaltbild (a) und Großsignal-Ersatzschaltbild (b).

$U_{BE} \approx 0,7 \text{ V}$ die nachfolgende Beziehung.

$$U_a \approx U_e - 0,7 \text{ V} \quad (3.70)$$

Damit folgt die Ausgangsspannung der Eingangsspannung mit einer Verschiebung von U_{BE} . Deshalb wird die Kollektorschaltung oft als Emitterfolger bezeichnet.

Kleinsignalverhalten der Kollektorschaltung

Wie im vorigen Abschnitt soll die Schaltung nur durch ein kleines Signal um den Arbeitspunkt A angesteuert werden. Dabei entspricht die Kleinsignalverstärkung der Steigung der Übertragungskennlinie. Durch Differentiation von Gl. 3.70 erhält man

$$A = \left. \frac{\partial U_a}{\partial U_e} \right|_A \approx 1 \quad (3.71)$$

Mit dem Kleinsignalersatzschaltbild nach Bild 3.41 und dem Knotensatz erhält man

$$\frac{u_e - u_E}{r_{BE}} + S \cdot u_{BE} = \left(\frac{1}{R_E} + \frac{1}{r_{CE}} \right) \cdot u_a \quad (3.72)$$

Nach Auflösung mit $u_{BE} = u_e - u_E$ folgt

$$A = \left. \frac{u_a}{u_e} \right|_{i_a=0} = \frac{\left(1 + \frac{1}{\beta}\right) \cdot S \cdot R'_E}{\left(1 + \frac{1}{\beta}\right) \cdot S \cdot R'_E + 1} \quad (3.73)$$

$$\underset{\beta \gg 1}{\approx} \underset{r_{CE} \gg R_E}{\frac{S \cdot R_E}{S \cdot R_E + 1}} \underset{S R_E \gg 1}{\approx} 1$$

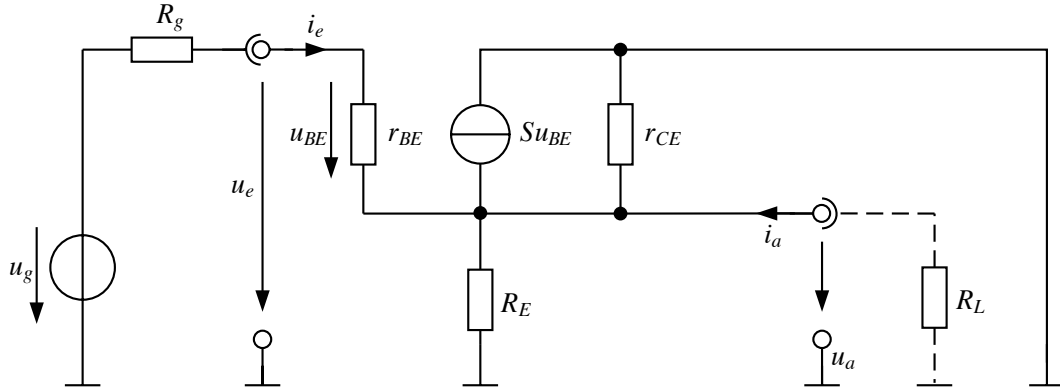


Bild 3.41: Kleinsignal-Ersatzschaltbild der Kollektorschaltung.

Der Kleinsignal-Eingangswiderstand beträgt

$$r_e = \left. \frac{u_e}{i_e} \right|_{i_a=0} = r_{BE} + (1 + \beta) \cdot R'_E \stackrel{\substack{r_{CE} \gg R_E \\ \beta \gg 1}}{=} r_{BE} + \beta \cdot R_E \stackrel{\beta R_E \gg r_{BE}}{\approx} \beta \cdot R_E \quad (3.74)$$

Der Einfluss des Lastwiderstandes R_L kann durch Einsetzen der Parallelschaltung von R_E und R_L erfolgen. Für den Kleinsignal-Ausgangswiderstand ergibt sich

$$r_a = \frac{u_a}{i_a} = R'_E \parallel \frac{R_g + r_{BE}}{1 + \beta} \stackrel{\substack{r_{CE} \gg R_E \\ \beta \gg 1}}{=} R_E \parallel \left(\frac{R_g}{\beta} + \frac{1}{S} \right) \quad (3.75)$$

Er hängt vom Innenwiderstand des Signalgenerators ab. Man kann drei Bereiche unterscheiden:

$$r_a = \begin{cases} \frac{1}{S} & \text{für } R_g < r_{BE} = \frac{\beta}{S} \\ \frac{R_g}{\beta} & \text{für } r_{BE} < R_g < \beta \cdot R_E \\ R_E & \text{für } R_g > \beta \cdot R_E \end{cases} \quad (3.76)$$

Im Bereich von $r_{BE} < R_g < R_E$ erfolgt eine Transformation des Innenwiderstandes aus $r_a \approx R_g/\beta$. Deshalb bezeichnet man die Kollektorschaltung häufig als Impedanzwandler. Damit kann man eine Signalquelle mit einer nachfolgenden Kollektorschaltung auch in Form einer äquivalenten Signalquelle beschreiben, wie es in Bild 3.42 dargestellt ist. Zusammenfassend erhält man für eine Kollektorschaltung für $r_{CE} \gg R_E$, $\beta \gg 1$ und $R_L = \infty$ folgende Zusammenhänge:

$$A = \left. \frac{u_a}{u_e} \right|_{i_a=0} \approx \frac{S \cdot R_E}{1 + S \cdot R_E} \stackrel{S R_E \gg R_E}{\approx} 1 \quad (3.77)$$

$$r_e = \frac{u_e}{i_e} \Big|_{i_a=0} \approx r_{BE} + \beta \cdot R_E \stackrel{S R_E \gg 1}{\approx} \beta \cdot R_E \quad (3.78)$$

$$r_a = \frac{u_a}{i_a} \approx R_E \parallel \left(\frac{R_g}{\beta} + \frac{1}{S} \right) \quad (3.79)$$

Da die Differenz zwischen der Ausgangs- und Eingangsspannung gerade U_{BE} ist, beträgt die Temperaturdrift der Ausgangsspannung 2 mV/K.

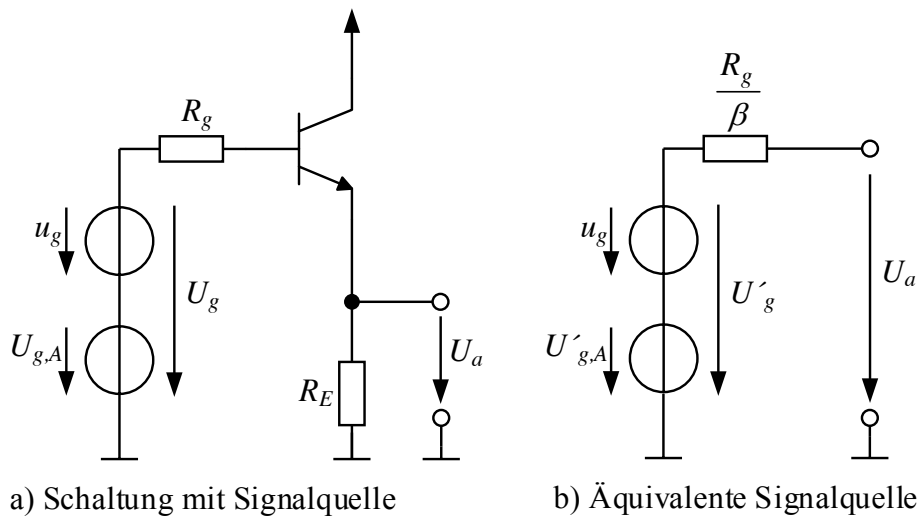


Bild 3.42: Kollektorschaltung als Impedanzwandler.

Arbeitspunkteinstellung

Die Einstellung eines stabilen Arbeitspunktes ist bei der Kollektorschaltung einfacher als bei der Emitterschaltung, da die Übertragungskennlinie über einen großen Bereich linear ist. Damit haben kleine Abweichungen praktische keine Auswirkungen auf das Kleinsignalverhalten. Man unterscheidet zwischen einer Wechselspannungskopplung und einer Gleichspannungskopplung. Bei der Wechselspannungskopplung werden die Signalquelle und Last über Koppelkondensatoren angeschlossen. Somit kann man den Gleichstromarbeitspunkt unabhängig von den Gleichspannungen der Signalquelle und der Last wählen. Nach Bild 3.43a wird die erforderliche Basisspannung durch den Spannungsteiler R_1 und R_2 eingestellt. Dabei wird der Querstrom durch die beiden Widerstände wesentlich größer als der Basisstrom gewählt. Somit ergibt sich die erforderliche Spannung zu

$$U_{B,A} = (I_{C,A} + I_{B,A}) \cdot R_E + U_{BE,A} \approx I_{C,A} \cdot R_E + 0,7 \text{ V} \quad (3.80)$$

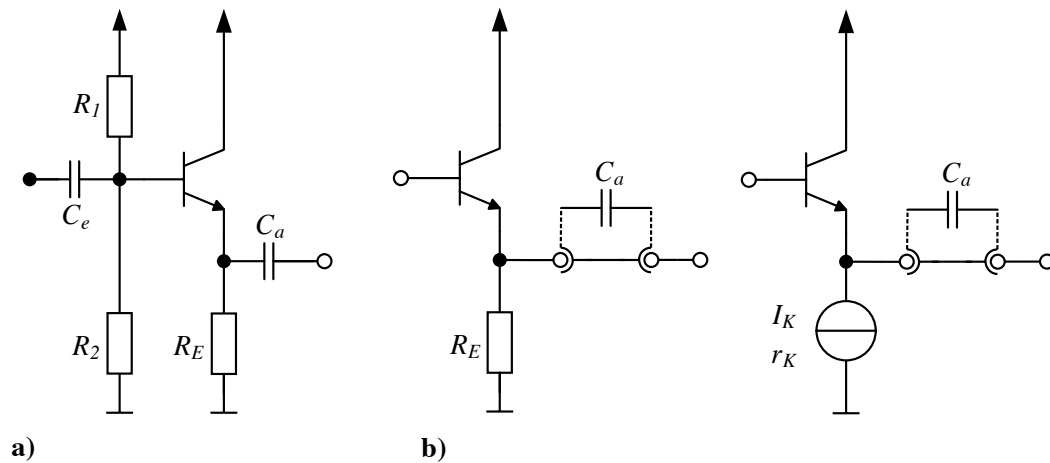


Bild 3.43: Arbeitspunkteinstellung der Kollektorschaltung.

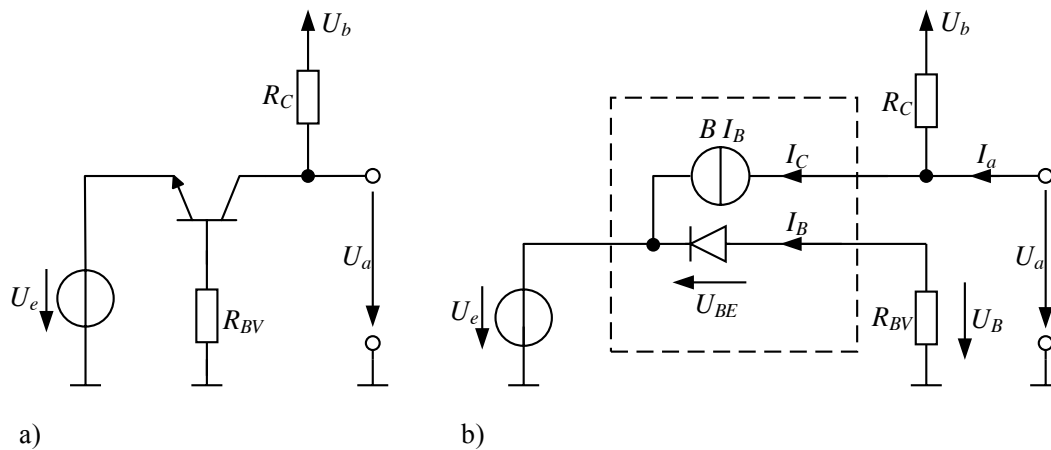


Bild 3.44: Basisschaltung (a) und Großsignal-Ersatzschaltbild (b).

In integrierten Schaltungen wird häufig eine Gleichspannungskopplung am Eingang verwendet. Dadurch entfallen die Widerstände R_1 und R_2 sowie der Koppelkondensator (siehe Bild 3.43b). In diesem Fall wird die Basisspannung durch die Ausgangsgleichspannung der vorangegangenen Stufe bestimmt. Die Koppelkondensatoren beeinflussen den Frequenzgang der Schaltung in vergleichbarer Weise wie im Fall der Emitterschaltung. Deshalb sind die Überlegungen aus Kapitel 3.2.4.1 gültig.

3.2.4.3 Die Basisschaltung

Bild 3.44a zeigt eine einfach aufgebaute Basisschaltung. Der Widerstand R_{BV} dient der Begrenzung des Basisstromes bei Übersteuerung. Im Normalbetrieb hat er faktisch keinen Einfluss. Damit die Kennlinien nicht vom Innenwiderstand R_g der Signalquelle abhängen, wurde eine Signalspannungsquelle ohne Innenwiderstand (ideale Spannungs-

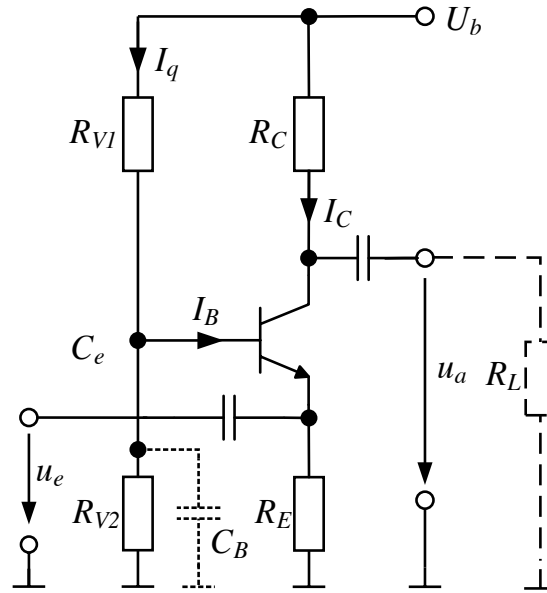


Bild 3.45: Komplett aufgebaute Basisschaltung für Wechselspannungskopplung mit Arbeitspunkteinstellung.

quelle) ausgewählt. Bild 3.44b zeigt das Großsignal-Ersatzschaltbild für den Normalbetrieb. Mit dem vereinfachten Transistormodell ergibt sich:

$$I_C = B \cdot I_B = B \cdot I_S \cdot e^{\frac{U_{BE}}{U_T}} \quad (3.81)$$

Aus Bild 3.44b folgt damit

$$U_a = U_b + (I_a - I_C) \cdot R_C \stackrel{I_a=0}{=} U_b - I_C \cdot R_C \quad (3.82)$$

$$U_e = -U_{BE} - I_B \cdot R_{BV} = -U_{BE} - \frac{I_C \cdot R_{BV}}{B} \approx -U_{BE} \quad (3.83)$$

Für $B \gg 1$ und kleine R_{BV} kann man den Spannungsabfall an R_{BV} vernachlässigen.

3.2.4.4 Kleinsignalverhalten der Basisschaltung

Die Berechnung der Kleinsignalparameter soll anhand einer komplett aufgebauten Basisschaltung nach Bild 3.45 mit Hilfe des Kleinsignalersatzschaltbildes nach Bild 3.46 erfolgen. Aus dem Knotensatz ergibt sich:

$$\frac{u_a}{R_C} + \frac{u_a - u_e}{r_{CE}} + S \cdot u_{BE} = 0 \quad (3.84)$$

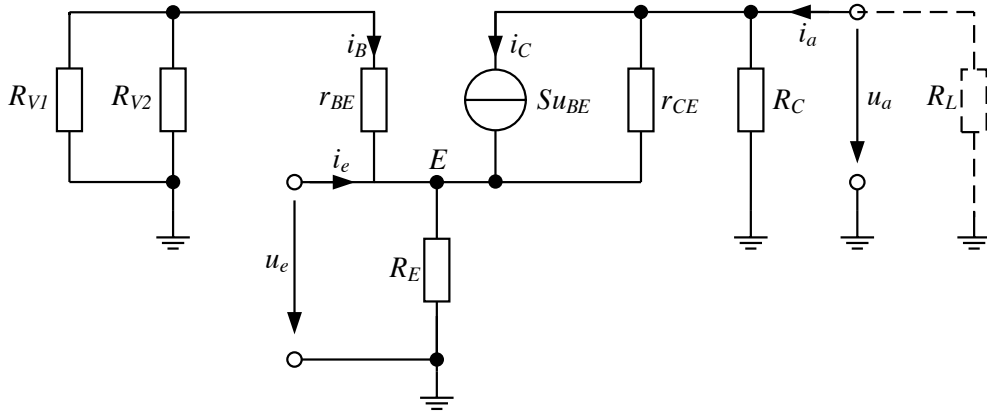


Bild 3.46: Kleinsignalersatzschaltbild der Basisschaltung nach Bild 3.45.

Mit der folgenden Spannungsteilung ergibt sich

$$u_{BE} = -\frac{r_{BE}}{r_{BE} + R_{BV}} \cdot u_e \quad (3.85)$$

mit $R_{BV} = R_{V1} \parallel R_{V2}$. Die Spannungsverstärkung A beträgt somit

$$A = \frac{u_a}{u_e} \bigg|_{i_a=0} = \left(\frac{\beta}{r_{BE} + R_{BV}} + \frac{1}{r_{CE}} \right) \cdot (R_C \parallel r_{CE}) \quad (3.86)$$

$r_{CE} \gg R_C$
 $\beta r_{CE} \gg r_{BE} + R_{BV} \approx \frac{\beta \cdot R_C}{r_{BE} + R_{BV}} \stackrel{r_{BE} \gg R_{BV}}{\approx} S \cdot R_C$

Die Verstärkung wird maximal bei $R_{BV} = 0$. Dies wird erreicht, durch das Hinzufügen des in Bild 3.45 gestrichelt eingezeichneten Kondensators C_B . Dadurch wird für Wechselspannungen die Parallelschaltung aus R_{V1} und R_{V2} quasi „kurzgeschlossen“. Der Kleinsignal-Eingangswiderstand der Schaltung beträgt

$$r_e = \frac{u_e}{i_e} \bigg|_{i_a=0} = R_E \parallel \left\| \frac{(r_{BE} + R_{BV})}{\beta} \right\| R_C + r_{CE} \quad (3.87)$$

$r_{CE} \gg R_C$
 $\stackrel{r_{CE} \gg R_C}{\approx} R_E \parallel \left(\frac{1}{S} + \frac{R_{BV}}{\beta} \right)$
 $\stackrel{r_{BE} \gg R_{BV}}{\approx} R_E \parallel \frac{1}{S} \stackrel{R_E \rightarrow \infty}{\approx} \frac{1}{S}$

Der Kleinsignal-Ausgangswiderstand beträgt

$$r_a = \frac{u_a}{i_a} = R_C \parallel \left\| r_{CE} + r_e \right\| \stackrel{r_{CE} + r_e \gg R_C}{\approx} R_C \quad (3.88)$$

Zusammenfassend ergeben sich mit $r_{CE} \gg R_C$, $\beta \cdot r_{CE} \gg r_{BE} + R_{BV}$, $\beta \gg 1$ folgende Gleichungen für die Basisschaltung:

$$A = \left. \frac{u_a}{u_e} \right|_{i_a=0} \approx \frac{\beta \cdot R_C}{r_{BE} + R_{BV}} \stackrel{r_{BE} \gg R_{BV}}{\approx} S \cdot R_C \quad (3.89)$$

$$r_e = \frac{u_e}{i_e} \approx R_E \left\| \left(\frac{1}{S} + \frac{R_{BV}}{\beta} \right) \stackrel{r_{BE} \gg R_{BV}}{\approx} R_E \right\| \frac{1}{S} \stackrel{R_E \rightarrow \infty}{\approx} \frac{1}{S} \quad (3.90)$$

$$r_a = \frac{u_a}{i_a} \approx R_C \quad (3.91)$$

Ein Vergleich mit den Gleichungen für die Emitterschaltung zeigt, dass das Kleinsignalverhalten der Basis- und Emitterschaltung sehr ähnlich ist. In beiden Schaltungen steuert der Signalgenerator die Basis-Emitterdiode und der Ausgang ist am Kollektor angeschlossen. Der einzige Unterschied besteht darin, dass der Eingangswiderstand um den Betrag β kleiner ist, weil der Emitterstrom $i_E \approx -\beta i_B$ an Stelle des Basisstroms als Eingangsstrom wirkt. Die Temperaturdrift verhält sich genauso wie bei der Emitterschaltung ohne Gegenkopplung:

$$\left. \frac{\partial U_a}{\partial T} \right|_A = \left. \frac{\partial U_a}{\partial U_e} \right|_A \cdot \frac{\partial U_e}{\partial T} \approx A \cdot 2 \frac{mV}{K} \quad (3.92)$$

Arbeitspunkteinstellung

Für die Basisschaltung nach Bild 3.45 wird die Spannung im Arbeitspunkt über den Basisspannungsteiler R_{V1} und R_{V2} eingestellt. Dabei wird der Querstrom I_q durch die beiden Widerstände wesentlich größer als der Basisstrom I_B gewählt.

3.3 Sperrschicht-Feldeffekttransistoren (JFET)

Lernziele

- Kennenlernen der Eigenschaften von Sperrschicht-Feldeffekttransistoren
- Funktion von Sperrschicht-Feldeffekttransistoren
- Arbeiten mit Eingangs- und Ausgangskennlinien
- Arbeitspunkteinstellung, Großsignalverhalten, Kleinsignalverhalten
- Grundschaltungen mit Sperrschicht-Feldeffekttransistoren

3.3.1 Aufbau und Wirkungsweise

Wie bei bipolaren Transistoren, so ist auch bei Sperrschicht-Feldeffekttransistoren ein pn-Übergang von Bedeutung. Im Gegensatz zur Bipolartechnik wird der steuernde pn-Übergang jedoch immer in Sperrrichtung betrieben. Die Anschlüsse des Feldeffekttransistors werden mit Source (Quelle), Drain (Senke) und Gate (Tor) bezeichnet. Der schematische Aufbau eines n- und eines p-Kanal Sperrschicht-Feldeffekttransistors und die zugehörigen Schaltungssymbole sind in Bild 3.47 dargestellt. Zur weiteren Erklärung

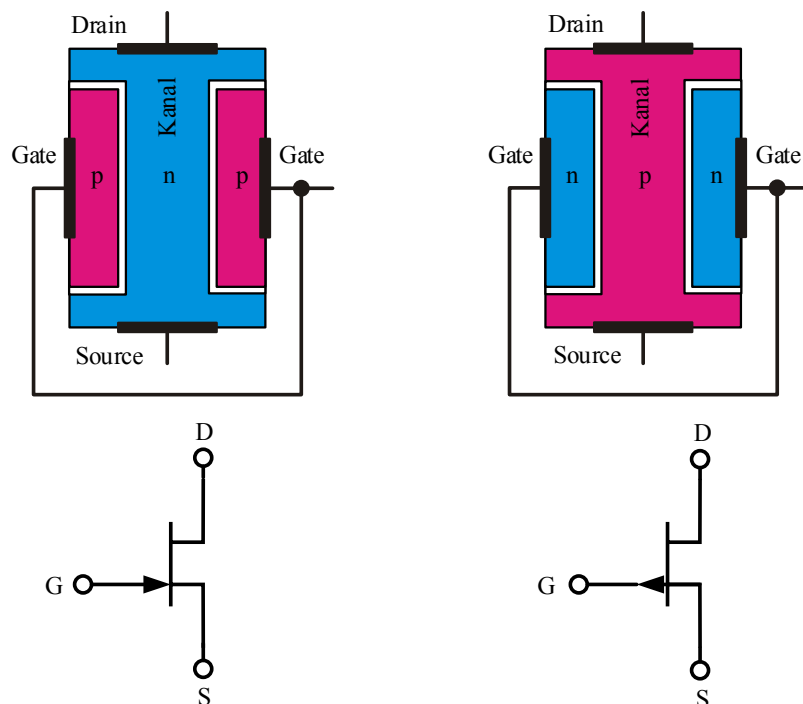


Bild 3.47: Schematische Darstellung und Schaltungssymbole von n- und p-Kanal Sperrschicht-Feldeffekttransistoren.

der Funktion soll hier nur der n-Kanal Transistor betrachtet werden. Zwischen den Anschlüssen S und D ist nur das n-dotierte Halbleitermaterial vorhanden. Dies wird als Kanal bezeichnet. Um die p-dotierten Gates bildet sich eine Raumladungszone (Sperrschicht), deren Gesamtbreite sich überwiegend im n-dotierten Kanal ausbreitet, da der Kanal wesentlich niedriger dotiert ist, als die Gatebereiche. Wird nun zwischen Drain und Source eine kleine positive Spannung angelegt ($U_{DS} \ll 1$ V) und das Gate mit dem Source verbunden, so dass $U_{GS} = 0$ V ist, so kann ein kleiner Strom zwischen den beiden Anschlüssen fließen. Man bezeichnet diese Art von Feldeffekttransistoren auch als selbstleitend. Die Stromstärke im Kanal wird bestimmt durch den ohm'schen Widerstand des Materials.

$$R = \rho \cdot \frac{\ell}{A} \quad \text{mit} \quad \rho = \frac{1}{n \cdot e \cdot \mu_{eff}} \quad \text{und} \quad A = t \cdot w \quad (3.93)$$

Die eingesetzten Größen bedeuten hierbei:

- n : Dichte der Ladungsträger im Kanal [cm^{-3}]
- e : Elementarladung [As]
- μ_{eff} : Beweglichkeit der Ladungsträger [cm/Vs]
- ℓ : Länge (length) des Kanals (hier Abstand zwischen Source und Drain)
- w : Breite (width) des Kanals
- t : Dicke (thickness) des Kanals (Abstand zwischen den Raumladungszonen)

Wird nun an das Gate eine negative Spannung angelegt, so wird die Raumladungszone breiter (Bild 3.48) und damit die Dicke t des Kanals kleiner. Als direkte Auswirkung wird der Widerstand des Kanals größer. Dies kann solange fortgesetzt werden, bis sich die Raumladungszonen berühren und der Kanal unterbrochen wird. Dann fließt kein Strom mehr zwischen Drain und Source. Man spricht dabei auch von einer Abschnürung des Kanals. Die dafür notwendige Spannung am Gate bezeichnet man als Abschnürspannung (pinch-off voltage bzw. threshold voltage). Die gebräuchlichsten Bezeichnungen dafür sind U_P bzw. U_{th} . Wir werden im Weiteren die Bezeichnung U_{th} verwenden.

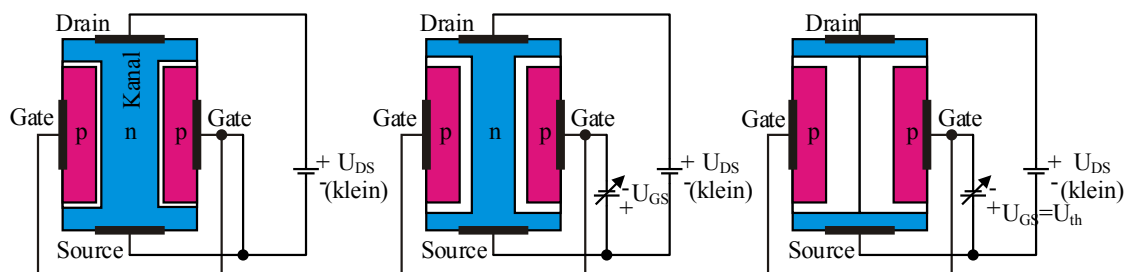


Bild 3.48: Ausdehnung der Raumladungszone beim JFET durch Verändern der Gate-Source-Spannung.

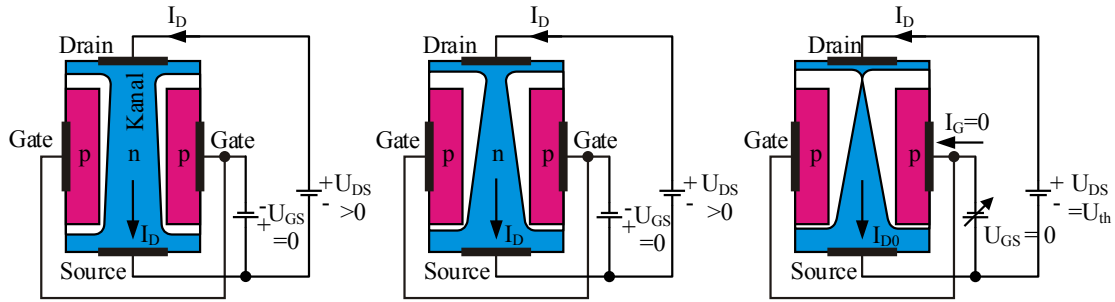


Bild 3.49: Ausdehnung der Raumladungszone beim JFET durch Verändern der Gate–Source–Spannung.

In dem bisher besprochenen Arbeitsbereich, also bei sehr kleinen Werten der Drain–Source–Spannung, verhält sich der Sperrschicht–Feldeffekttransistor wie ein steuerbarer ohm’scher Widerstand. In der analogen Schaltungstechnik werden aber Versorgungsspannungen zwischen $\pm 5\text{ V}$ und $\pm 15\text{ V}$ eingesetzt. Daraus ergibt sich ein anderes Verhalten des Sperrschicht–Feldeffekttransistors als bisher beschrieben. Um dies zu erläutern wird die Gate–Source–Spannung $U_{GS} = 0\text{ V}$ angelegt und die Drain–Source–Spannung von 0 V beginnend kontinuierlich erhöht. Durch den ohm’schen Widerstand des Kanalmaterials nimmt das Potential vom Source (der mit Masse verbunden wird) zum Drain zu. Damit steigt die im Kanal wirksame Sperrspannung in gleicher Weise, was zu einer Verbreiterung der Raumladungszone in der Nähe des Drain führt (Bild 3.49), bis sich schließlich die linke und die rechte Raumladungszone berühren und dadurch der Kanal an einer Stelle abgeschnürt wird. Dies geschieht bei der Spannung $U_{DS} = -U_{th}$. Trotz der Abschnürung des Kanals fließt jetzt aber ein Strom $I_D = I_{D0}$. Eine weitere Erhöhung von U_{DS} führt nur noch zu einer geringen Erhöhung des Drainstroms. In der Literatur wird I_{D0} auch gerne mit I_{DSS} bezeichnet. Da die einwandfreie Funktion eines Sperrschicht–Feldeffekttransistors von der Existenz eines gesperrten pn–Übergangs zwischen Gate und Kanal abhängt, muss darauf geachtet werden, dass beim n–Kanal JFET nur negative Gate–Sourcespannungen angelegt werden.

3.3.2 Funktionen eines Sperrschicht–Feldeffekttransistors

3.3.2.1 Eingangs– und Ausgangskennlinien

Die Eingangskennlinie und das Ausgangskennlinienfeld eines idealisierten Sperrschicht–Feldeffekttransistors sind in Bild 3.50 dargestellt. Die Gleichung der Eingangskennlinie ist:

$$I_D = I_{D0} \cdot \left(1 - \frac{U_{GS}}{U_{th}}\right)^2 \quad (3.94)$$

Für $U_{GS} = 0\text{ V}$ wird damit der Drainstrom zu $I_D = I_{D0}$ und für $U_{GS} = U_{th}$ wird $I_D = 0$. Das Ausgangskennlinienfeld lässt sich durch drei Bereiche beschreiben:

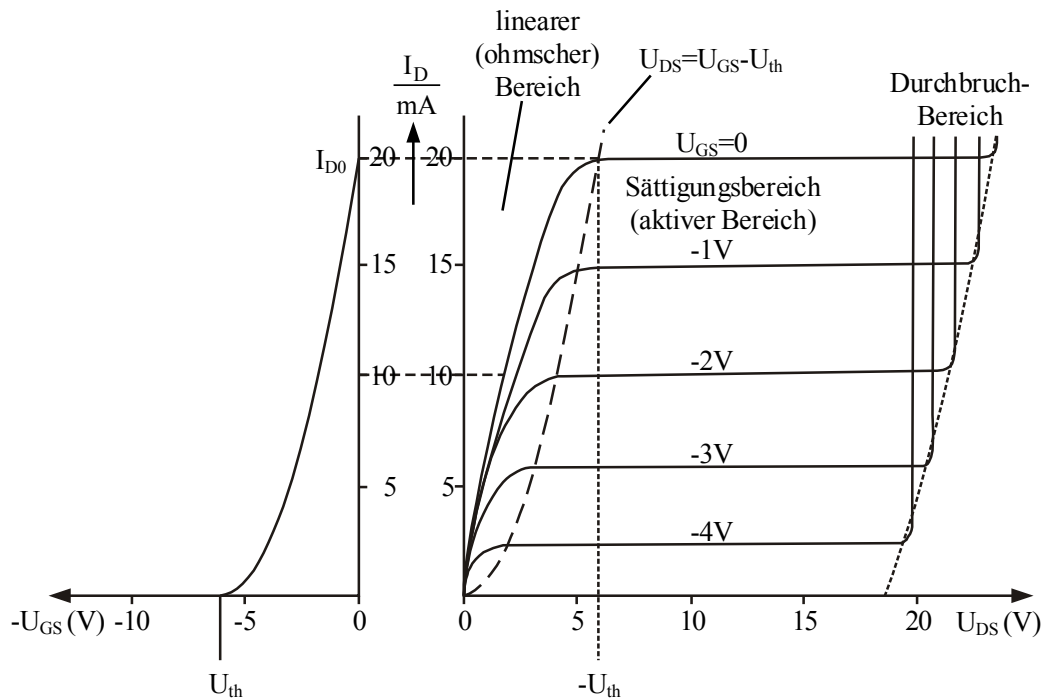


Bild 3.50: Eingangskennlinie und Ausgangskennlinienfeld.

- den linearen oder ohm'schen Bereich.
- den Sättigungsbereich (Abschnürbereich, aktiver Bereich), der für analoge Schaltungen der Arbeitsbereich ist und
- den Durchbruchbereich

Im linearen Bereich gilt:

$$I_D = I_{D0} \cdot \left[2 \cdot \left(1 - \frac{U_{GS}}{U_{th}} \right) \cdot \frac{U_{DS}}{(-U_{th})} - \left(\frac{U_{DS}}{U_{th}} \right)^2 \right] \quad (3.95)$$

Diese Darstellungsform ist die in der Literatur gebräuchlichste. Durch Umformung entsteht die folgende Gleichung:

$$I_D = \frac{I_{D0}}{U_{th}^2} \cdot [2 \cdot (U_{GS} - U_{th}) \cdot U_{DS} - (U_{DS})^2] \quad (3.96)$$

Bringt man noch den Faktor 2 vor die Klammer, erhält man

$$I_D = 2 \cdot \frac{I_{D0}}{U_{th}^2} \cdot \left[(U_{GS} - U_{th}) \cdot U_{DS} - \frac{(U_{DS})^2}{2} \right] \text{ mit } 2 \cdot \frac{I_{D0}}{U_{th}^2} = \beta \quad (3.97)$$

Mit dieser Form der Gleichung soll zukünftig weitergearbeitet werden. Der Faktor β ist

beim Feldeffekttransistor im Gegensatz zum bipolaren Transistor kein dimensionsloser Stromverstärkungsfaktor, sondern der so genannte Steilheitskoeffizient mit der Dimension $[A/V^2]$. Setzt man in die vorherige Gleichung den Wert von U_{DS} an der Grenze zwischen linearem und Sättigungsbereich $U_{DS} = (U_{GS} - U_{th})$ ein (gestrichelte Linie im Kennlinienfeld), wird

$$I_D = \frac{I_{D0}}{U_{th}^2} \cdot (U_{GS} - U_{th})^2 = \frac{1}{2} \cdot \beta (U_{GS} - U_{th})^2 \quad (3.98)$$

Eine weitere wichtige Kenngröße für analoge Anwendungen eines Transistors ist die Steilheit S , die bei Feldeffekttransistoren auch oft mit g_m bezeichnet wird. Sie ist definiert als:

$$S = \frac{\partial I_D}{\partial U_{GS}} = 2 \cdot \frac{I_{D0}}{U_{th}^2} \cdot (U_{GS} - U_{th}) = \beta \cdot (U_{GS} - U_{th}) \quad (3.99)$$

Betrachten wir nun noch den differentiellen Widerstand r_{DS} des Transistors zwischen dem Source- und dem Drain-Anschluss. Im linearen oder ohm'schen Bereich und für kleine Werte von U_{DS} kann I_D näherungsweise durch

$$I_D = 2 \cdot \frac{I_{D0}}{U_{th}^2} \cdot (U_{GS} - U_{th}) \cdot U_{DS} \quad (3.100)$$

beschrieben werden. Dann wird der differentielle Widerstand r_{DS} :

$$r_{DS} = \frac{\partial U_{DS}}{\partial I_D} = \frac{U_{th}^2}{2 \cdot I_{D0} \cdot (U_{GS} - U_{th})} \quad (3.101)$$

Im Sättigungsbereich ist beim idealisierten Kennlinienfeld der Drainstrom konstant. Damit wird der differentielle Widerstand $r_{DS} \rightarrow \infty$.

3.3.2.2 Der Early-Effekt

Wie beim Bipolartransistor der Kollektorstrom, so steigt im Sättigungsbereich des Feldeffekttransistors auch der Drainstrom bei zunehmender Drain-Source-Spannung weiter an (Bild 3.51). Dieser Early-Effekt wurde bereits beim Bipolartransistor erläutert. Durch den Early-Effekt verändern sich die Gleichungen für den Drainstrom und die Steilheit zu:

$$\begin{aligned} I_D &= \beta \cdot \left[(U_{GS} - U_{th}) \cdot U_{DS} - \frac{(U_{DS})^2}{2} \right] \cdot \left(1 + \frac{U_{DS}}{U_A} \right) \\ I_D &= \frac{1}{2} \cdot \beta (U_{GS} - U_{th})^2 \cdot \left(1 + \frac{U_{DS}}{U_A} \right) \\ S &= \frac{\partial I_D}{\partial U_{GS}} = \beta \cdot (U_{GS} - U_{th}) \cdot \left(1 + \frac{U_{DS}}{U_A} \right) \end{aligned} \quad (3.102)$$

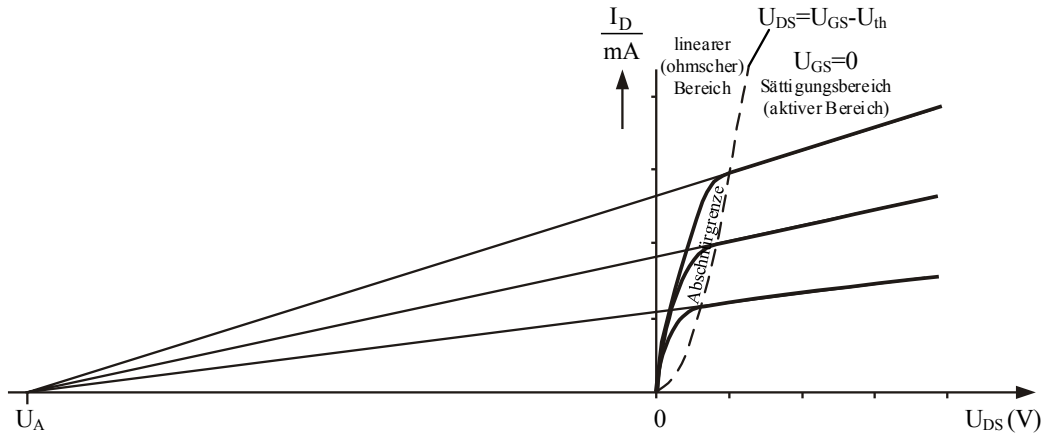


Bild 3.51: Eingangskennlinie und Ausgangskennlinienfeld.

Ist die Early-Spannung U_A bekannt, lässt sich der Drain-Sourcewiderstand durch folgende Gleichung bestimmen:

$$r_{DS} = \frac{\Delta U_{DS}}{\Delta I_D} = \frac{|U_A| + (U_{GS} - U_{th})}{I_D} = \frac{|U_A| + (U_{GS} - U_{th})}{\frac{I_{D0}}{U_{th}^2} \cdot (U_{GS} - U_{th})^2} \quad (3.103)$$

Damit nimmt r_{DS} nun einen endlich großen Wert an. Bei Sperrschicht-Feldeffekttransistoren liegt die Early-Spannung bei etwa einigen hundert Volt, d.h. die Kennlinien verlaufen im Arbeitsbereich nahezu waagrecht und der differentielle Drain-Sourcewiderstand r_{DS} ist noch immer sehr groß.

3.3.2.3 Bestimmung der Schwellspannung U_{th}

Häufig stellt sich für den Schaltungsentwickler die Frage nach dem genauen Wert der Schwellspannung U_{th} des in der Schaltung eingesetzten Transistors. Aus dem Kennlinienfeld des Transistors lassen sich im Sättigungsbereich Wertepaare U_{GS} und I_D (für $U_{DS} = \text{const.}$) ablesen. Aus der Gleichung des Drainstroms im Sättigungsbereich (Gl. 3.98) erhält man:

$$\sqrt{I_D} = \sqrt{\frac{1}{2} \cdot \beta} \cdot (U_{GS} - U_{th}) \quad (3.104)$$

Dies entspricht einer Geradengleichung der Form $y = m(x - x_0)$, so dass man durch die Erstellung eines Diagramms mit Wertepaaren $\sqrt{I_D}$, U_{GS} und graphischer Extrapolation die Schwellspannung als Schnittpunkt der so erhaltenen Geraden mit der Abszisse erhält, wie dies in Bild 3.52 gezeigt ist.

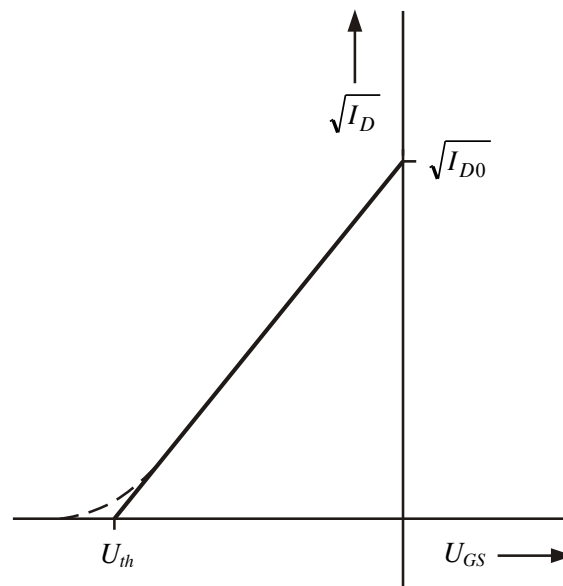


Bild 3.52: Graphische Ermittlung der Schwellspannung U_{th} eines JFET.

3.3.3 Arbeitspunkt und Kleinsignalverhalten

Wie bei bipolaren Transistoren, so ist auch beim Sperrschicht-Feldeffekttransistor die Verstärkung von sehr kleinen Signalen ein sehr wichtiges Anwendungsgebiet. Auch hier wird der JFET nur in einem kleinen Bereich um den Arbeitspunkt A angesteuert.

Im Gegensatz zum bipolaren Transistor, der spannungs- oder stromgesteuert betrieben werden kann, wird der Feldeffekttransistor nur spannungsgesteuert betrieben. Der Eingangsstrom des gesperrten pn-Übergangs I_G ist durch die Gleichung des Diodenstroms beschreibbar.

$$I_G = I_S \cdot \left(e^{\frac{U_{GS}}{U_T}} - 1 \right) \quad (3.105)$$

Da U_{GS} immer negativ ist, ist der Gatestrom beim JFET sehr klein ($10^{-9} - 10^{-12}$ A). Die Eingangs- oder Übertragungskennlinie eines JFET ist in Bild 3.53 dargestellt. Für zwei Eingangsspannungen ($U_{GS} = -1$ V und $U_{GS} = -3$ V) sind die Arbeitspunkte A_1 und A_2 eingetragen. Die Änderung des Drainstroms bei gleicher Änderung der Gate-Source-Spannung um $\pm 0,5$ V zeigt die Abhängigkeit der Stromänderung von der Wahl des Arbeitspunktes. Legt man die Tangente an die Arbeitspunkte, so erkennt man, dass die Steigung, die der Steilheit der Kennlinie im Arbeitspunkt entspricht, unterschiedlich ist. Die Ursache hierfür ist die Abhängigkeit des Drainstroms von $(U_{GS} - U_{th})^2$. Dies bedeutet aber auch, dass eine lineare Verstärkung des Eingangssignals nur in einem sehr kleinen Bereich um den sorgfältig ausgewählten Arbeitspunkt möglich ist. Das zugehörige Ausgangskennlinienfeld ist in Bild 3.54 gezeigt. Die beiden Arbeitspunkte A_1 und A_2 aus Bild 3.53 sind an den Schnittpunkten der entsprechenden Ausgangskennlinie mit einem angenommenen Lastwiderstand eingetragen. Wenn man die beiden Arbeitspunkte

betrachtet, erkennt man, dass beide nicht besonders günstig gewählt sind. Besser wäre sicher ein Arbeitspunkt um $U_{GS} = -2$ V, denn damit würde man etwa bei der Hälfte der für das Beispiel gewählten Versorgungsspannung von 20 V liegen.

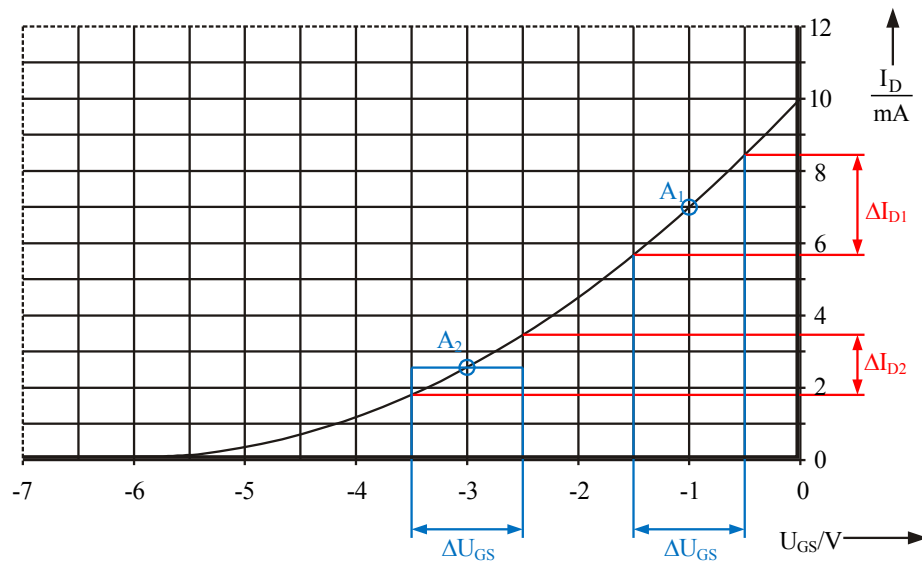


Bild 3.53: Eingangs- oder Übertragungskennlinie eines JFET.

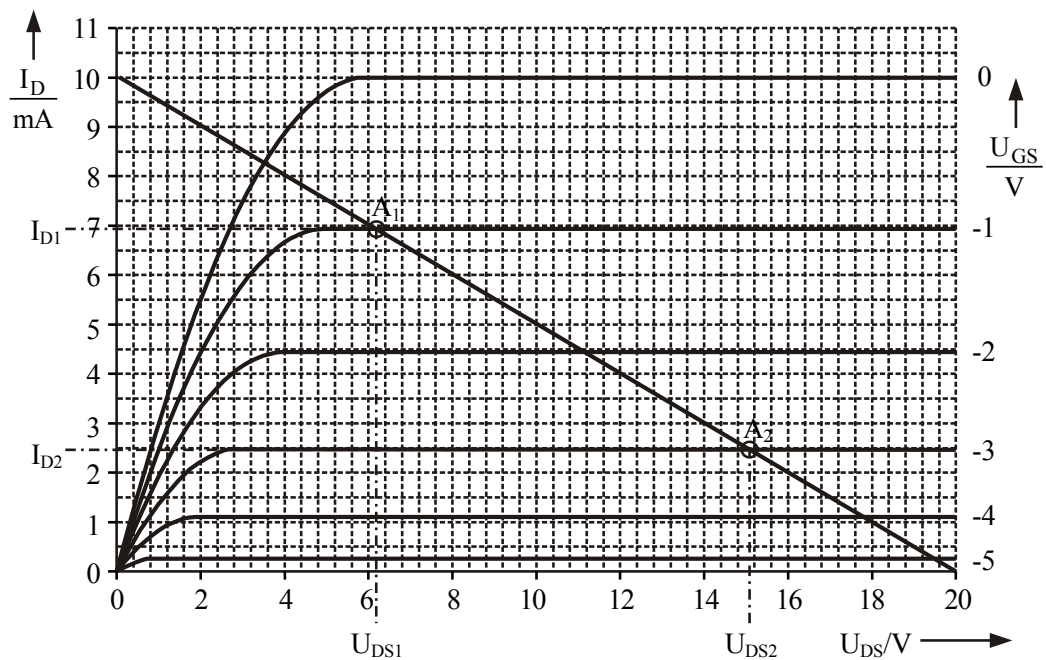


Bild 3.54: Ausgangskennlinienfeld des JFET mit zwei Arbeitspunkten.

3.3.3.1 Kleinsignal-Ersatzschaltbild

Das Kleinsignal-Ersatzschaltbild des JFET ist in Bild 3.55 gezeigt. Zwischen Gate und Drain bzw. Gate und Source liegen jeweils eine Diode und ein Kondensator parallel. Die Gate-Sourcekapazität C_{GS} ist als Eingangskapazität des Transistors zu interpretieren, während die Gate-Drainkapazität als eine parasitäre Kapazität (Rückkopplungskapazität) zwischen dem Ausgang und dem Eingang des Transistors gesehen werden kann. Dies wird vor allem bei hochfrequenten Signalen von Bedeutung. Da beide Kapazitäten, C_{GS} und C_{GD} , Sperrschichtkapazitäten sind, variiert der Kapazitätswert mit der Größe der anliegenden Spannungen. Die Widerstände R_S und R_D stellen den Drain- bzw. Sourcewiderstand des Halbleitermaterials dar. Dieses Modell wird für die Simulation von Schaltungen in PSPICE verwendet.

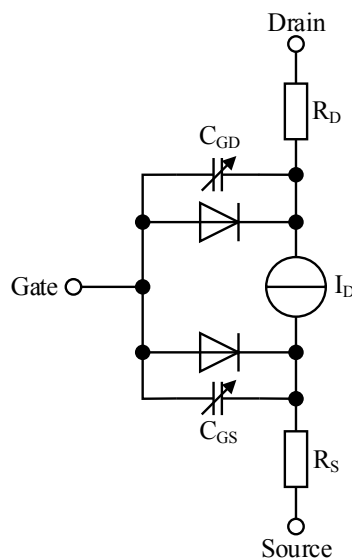


Bild 3.55: PSPICE-Ersatzschaltbild eines JFET.

Die Gleichungen, mit denen gearbeitet wird, entsprechen den bisher in diesem Kapitel angegebenen Gleichungen. Hinzu kommen jedoch noch eine Reihe von weiteren Effekten, wie z.B. der Temperatureffekt, und Rauscheffekte, um ein realistisches Verhalten des Transistors bei der Simulation zu berücksichtigen. Zur Übersicht und schnellen Berechnung von Schaltungen mit Sperrschicht-Feldeffekttransistoren wollen wir aber vereinfachte Ersatzschaltbilder, auf die im Folgenden näher eingegangen wird, verwenden. Die gleichstrommäßige Einstellung des Arbeitspunktes setzt voraus, dass nur negative Gate-Sourcespannungen angelegt werden, damit der pn-Übergang vom Gate zum Source immer in Sperrrichtung betrieben wird. Dazu muss entweder das Sourcepotential gegenüber dem Gatepotential durch den Einsatz eines Widerstands angehoben werden, oder mit zusätzlichen negativen Spannungsquellen und einem Spannungsteiler die richtige Gate-Sourcespannung eingestellt werden. Da in das Gate praktisch kein Eingangsgleichstrom fließt, ist ein solcher Spannungsteiler praktisch unbelastet. Ist der Arbeitspunkt gleichstrommäßig eingestellt, kann der Transistor für die Ansteuerung mit

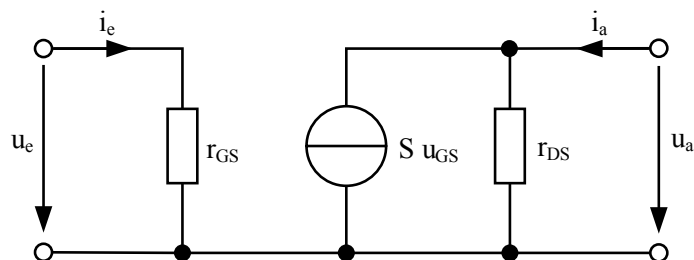


Bild 3.56: Kleinsignal-Ersatzschaltbild für Ansteuerung mit kleinen Wechselspannungen.

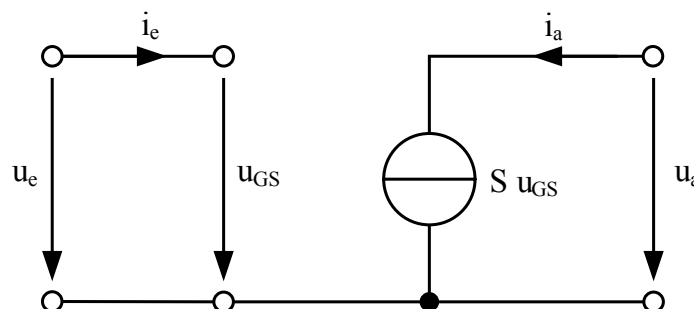


Bild 3.57: Vereinfachtes Kleinsignal-Ersatzschaltbild.

kleinen Wechselspannungen betrachtet werden. Bild 3.56 zeigt eine Variante für eine Wechselstromansteuerung. Der Widerstand r_{GS} ist der Eingangswiderstand, der Werte von $> 10^9 \Omega$ annehmen kann. Der Widerstand r_{DS} ist der differentielle Widerstand der Kennlinie im Sättigungsbereich. Beim idealisierten JFET können wir $r_{DS} \rightarrow \infty$ annehmen. Mit diesen Vereinfachungen ergibt sich das Ersatzschaltbild nach Bild 3.57. Der Eingangsstrom i_e wird zu Null und der Strom i_a entspricht dem Drainstrom.

3.3.3.2 Grundsaltungen mit Sperrschicht-Feldeffekttransistoren

Die drei Grundsaltungen mit Sperrschicht-Feldeffekttransistoren sind die Source-, die Drain- und die Gateschaltung. Wie bei den Grundsaltungen mit bipolaren Transistoren, sollen auch hier die Eigenschaften der Grundsaltungen wie Spannungsverstärkung, Eingangswiderstand und falls möglich der Ausgangswiderstand überschlägig berechnet werden. Die Sourceschaltung (Bild 3.58) hat als gemeinsamen Anschluss für Eingang und Ausgang den Source-Anschluss des Transistors. Das Eingangssignal wird am Gate angelegt und am Drain wird die Ausgangsspannung abgegriffen. Zur Bestimmung der wichtigen Größen ist das Wechselstromersatzschaltbild nach Bild 3.59 hilfreich. Beginnen wir mit der Ermittlung des Eingangswiderstands r_e der Sourceschaltung:

$$r_e = \frac{u_e}{i_e} = r_{GS} \quad \Rightarrow r_e \text{ sehr groß} \quad (3.106)$$

Die Spannungsverstärkung der Schaltung ist:

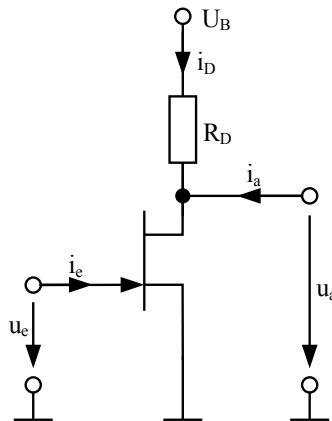


Bild 3.58: Die Sourceschaltung.

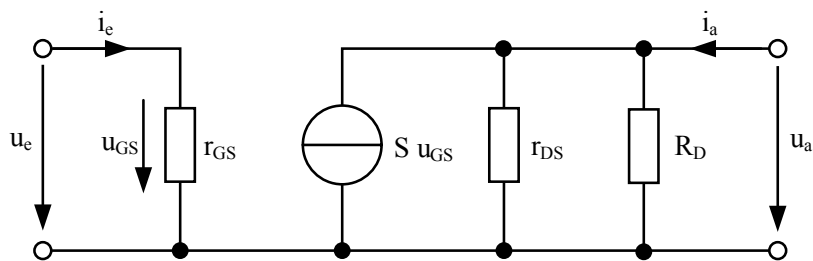


Bild 3.59: Kleinsignal-Ersatzschaltbild der Sourceschaltung.

$$A = \frac{u_a}{u_e}, \quad u_e = u_{GS}, \quad u_a = -S \cdot u_{GS} (R_D \parallel r_{DS}) \quad (3.107)$$

$$A = -S \cdot (R_D \parallel r_{DS}) \approx -S \cdot R_D \quad \text{da} \quad r_{DS} \gg R_D$$

Diese Gleichung kennen wir schon von der Emitterschaltung mit einem bipolaren npn Transistor. Beim Feldeffekttransistor ist jedoch die Steilheit etwa um den Faktor 10 kleiner, also wird auch die Spannungsverstärkung um diesen Faktor kleiner. Der Ausgangswiderstand r_a wird zu

$$r_a = \frac{u_a}{i_a} = (R_D \parallel r_{DS}) \approx R_D \quad (3.108)$$

Daraus ist zu erntnehmen, dass der Ausgangswiderstand wesentlich durch die externe Beschaltung des Transistors bestimmt wird.

Die Drainschaltung in Bild 3.60 wird in der Literatur oft als Sourcefolger bezeichnet. Zur Bestimmung der wichtigen Größen ist auch hier das Wechselstromersatzschaltbild (Bild 3.61) hilfreich. Beginnen wir zunächst mit der Spannungsverstärkung der Drainschaltung:

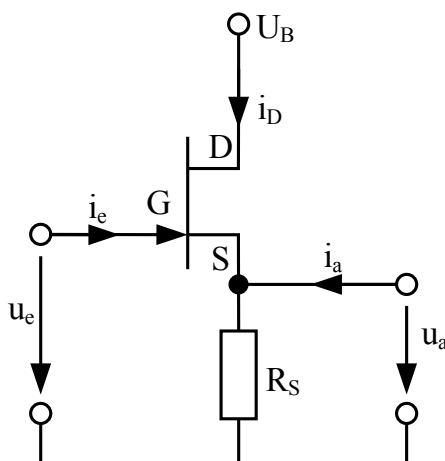


Bild 3.60: Die Drainschaltung.

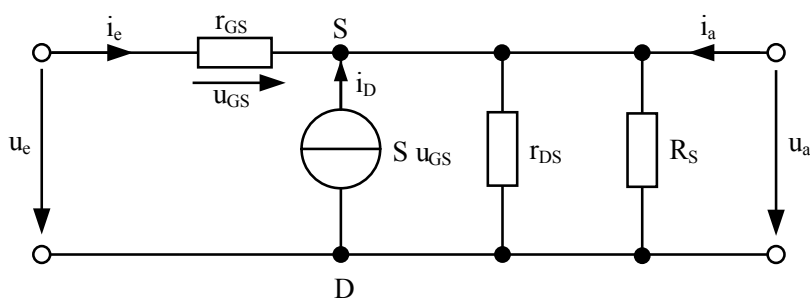


Bild 3.61: Kleinsignal-Ersatzschaltbild der Drainschaltung.

$$\begin{aligned}
 A = \frac{u_a}{u_e} \quad \text{mit} \quad r_{DS} \gg R_S \quad \text{ist} \quad i_a \approx i_D \\
 u_e = u_{GS} + u_a \quad u_{GS} = r_{GS} \cdot i_e \\
 i_D = S \cdot u_{GS} \\
 u_a = i_D \cdot R_S \Rightarrow u_a = S \cdot u_{GS} \cdot R_S \\
 A = \frac{S \cdot u_{GS} \cdot R_S}{u_{GS} + S \cdot u_{GS} \cdot R_S} = \frac{S \cdot R_S}{1 + S \cdot R_S} = \frac{1}{1 + \frac{1}{S \cdot R_S}} < 1
 \end{aligned} \tag{3.109}$$

Der Eingangswiderstand r_e der Schaltung wird zu:

$$r_e = \frac{u_e}{i_e} = \frac{r_{GS} \cdot u_e}{u_{GS}} = \frac{r_{GS} \cdot u_{GS} \cdot (1 + S \cdot R_S)}{u_{GS}} = r_{GS} \cdot (1 + S \cdot R_S) \tag{3.110}$$

Nehmen wir $S = 5 \text{ mS}$ und $R_S = 2 \text{ k}\Omega$ an, wird $S \cdot R_S = 10$. Der Eingangswiderstand ist damit um den Faktor 11 größer, als bei der Sourceschaltung. Der Ausgangswiderstand

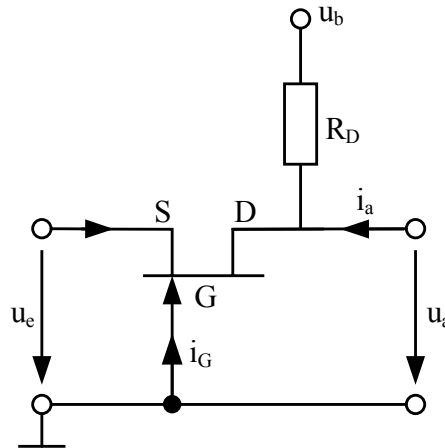


Bild 3.62: Die Gateschaltung.

r_a der Schaltung ergibt sich zu

$$r_a = \frac{u_a}{i_a} \quad \text{mit} \quad u_e = 0, \quad i_a = \frac{u_a}{R_S \parallel r_{DS} \parallel r_{GS}} - S \cdot u_{GS}$$

wird $i_a \approx u_a \cdot \left(\frac{1}{R_S} + S \right)$ für $u_{GS} = -u_a$, $r_{GS} \gg r_{DS} \gg R_S$ (3.111)

$$r_a = \frac{u_a}{u_a \cdot \left(\frac{1}{R_S} + S \right)} = \frac{R_S}{1 + S \cdot R_S} = \frac{\frac{1}{S} \cdot R_S}{\frac{1}{S} + R_S} = \frac{1}{S} \parallel R_S$$

Als letzte der drei Grundschaltungen betrachten wir die Gateschaltung nach Bild 3.62. Da die Gateschaltung nur in ganz seltenen Fällen in der Hochfrequenztechnik eingesetzt wird, soll hier nur auf die Spannungsverstärkung und den Eingangswiderstand eingegangen werden. Der Ausgangswiderstand ist bei dieser Schaltung praktisch ausschließlich von der Ausgangsbeschaltung abhängig. Die Spannungsverstärkung ist

$$A = \frac{u_a}{u_e} = \frac{i_a \cdot R_D}{u_e} \quad \text{mit} \quad i_a = -i_D, \quad i_D = S \cdot u_{GS} \quad \text{und} \quad u_e = -u_{GS} \quad \text{ist} \quad (3.112)$$

$$A = \frac{-S \cdot u_{GS} \cdot R_D}{-u_{GS}} = S \cdot R_D$$

Der Eingangswiderstand berechnet sich zu

$$r_e = \frac{u_e}{i_e} = \frac{u_e}{-i_D} = \frac{u_e}{S \cdot u_e} = \frac{1}{S} \quad (3.113)$$

Der Eingangswiderstand der Gateschaltung ist formelmäßig gleich dem Eingangswiderstand der Basisschaltung. Da wie schon mehrfach erwähnt die Steilheit beim JFET aber etwa um den Faktor 10 kleiner ist als die beim bipolaren Transistor, wird der Eingangswiderstand nun um diesen Faktor größer.

3.4 Isolierschicht–Feldeffekttransistoren (MOSFET)

Lernziele

- Kennenlernen der Eigenschaften von MOS–Feldeffekttransistoren
- Funktion von MOS–Feldeffekttransistoren
- Arbeiten mit Eingangs– und Ausgangskennlinien
- Arbeitspunkteinstellung, Großsignalverhalten, Kleinsignalverhalten
- Grundsaltungen mit MOS–Feldeffekttransistoren

3.4.1 Aufbau und Wirkungsweise

Im letzten Kapitel wurde der Sperrschicht–Feldeffekt–Transistor (JFET) besprochen. Beim JFET ist das Gate durch einen gesperrten pn–Übergang vom Kanal isoliert, und der Strom im Source–Drain–Kanal wird über die Größe der Raumladungszone zwischen Gate und Kanal gesteuert. Im Gegensatz zum JFET wird beim MOSFET das Gate durch eine Isolationsschicht vom Kanal getrennt. Darauf wird auch der Name zurückgeführt, Metal–Oxid–Semiconductor–Field–Effect Transistor, abgekürzt MOSFET.

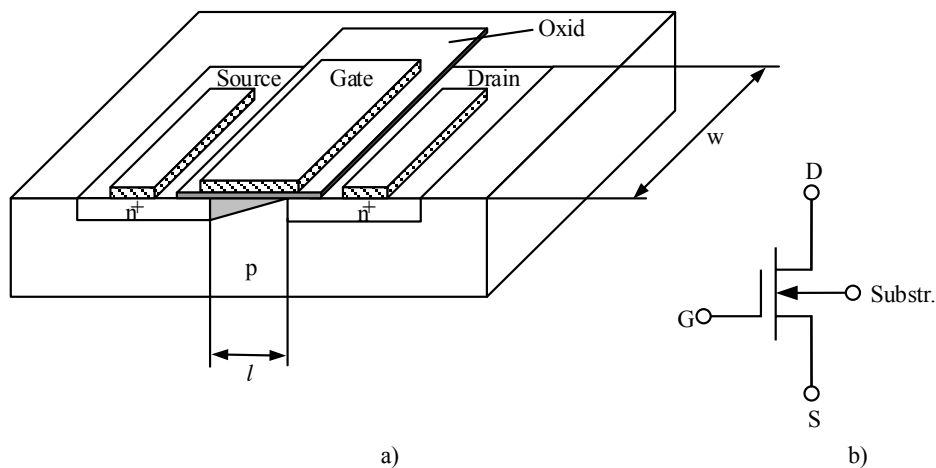


Bild 3.63: Schnittdarstellung des prinzipiellen Aufbaus eines MOSFET. (a) und Schaltungssymbol (b).

Bild 3.63a zeigt den Querschnitt eines MOSFET mit Schaltsymbol und Bild 3.64 schematisch den Aufbau eines n-Kanal Isolierschicht–Feldeffekttransistors. In einen p-dotierten Siliziumkristall werden zwei hoch dotierte n–Bereiche (Source und Drain) eingebracht. Der kleine Bereich zwischen Source und Drain wird als Kanal bezeichnet. Über der Kanalzone wird nun eine sehr dünne isolierende Oxidschicht aufgebracht, die wie

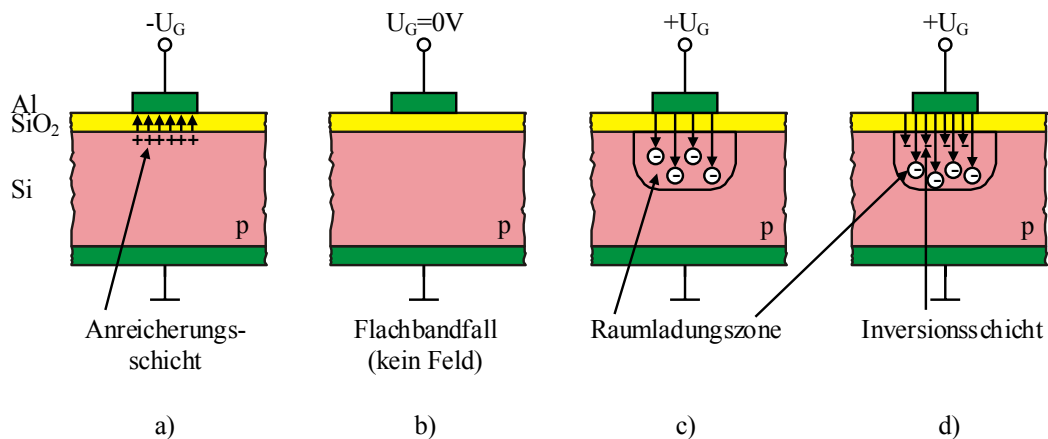


Bild 3.64: Schematische Darstellung der Ladungsverteilung in einer MOS-Struktur bei unterschiedlicher Spannung.

in Bild 3.64c dargestellt, noch einen Teil des Source- und des Drainbereichs mit überdeckt. Auf diese Isolierschicht wird jetzt das Gate aufgebracht, das entweder aus Metall, Silizium oder Metall-Siliziumverbindungen bestehen kann. Wichtig ist hierbei, dass der ohm'sche Widerstand des Materials möglichst klein ist. Bild 3.63b zeigt das Schaltungssymbol eines MOSFET. Zu den bereits bekannten Anschlüssen von Source, Drain und Gate kommt noch der Substratanschluss S dazu. Dieser Anschluss wird auch häufig als „bulk“ mit dem Symbol B bezeichnet. Das ist der Anschluss des Si-Trägersubstrates. Er wird bei diskreten MOSFETs in der Regel intern mit dem Sourceanschluss verbunden. In diesem Fall hat der MOSFET auch nur drei Anschlüsse. Bei genauer Betrachtung des Transistoraufbaus stellt man schnell fest, dass zum n-Kanal JFET ein wesentlicher Unterschied besteht: Die Dotierung der einzelnen Bereiche (Source, Drain und Kanal) ist beim MOSFET genau umgekehrt. Die Steuerspannung (Gatespannung) beeinflusst die Ladungsträgerdichte im Silizium unter dem Oxid. Je nach Polarität der Spannung und Dotierung des Siliziums bildet sich zwischen Source und Drain ein leitfähiger Kanal heraus, der n- oder p-leitend sein kann. Damit ist es durch die Dotierung möglich, selbstleitende oder selbstsperrende MOSFETs herzustellen; d.h. ein MOSFET ist selbstleitend, wenn bei einer kleinen Source-Drainspannung und bei $U_{GS} = 0 \text{ V}$ ein Strom fließt. Entsprechend ist ein MOSFET selbstsperrend, wenn bei einer kleinen Source-Drainspannung und bei $U_{GS} = 0 \text{ V}$ kein Strom fließt. Für analoge Schaltungen sind besonders die selbstleitenden Transistoren von Bedeutung, während in der Digitaltechnik die selbstsperrenden Transistoren eingesetzt werden. Zum besseren Verständnis der MOSFET soll zuerst nur die MOS-Struktur diskutiert werden, die in Bild 3.64 schematisch dargestellt ist. Diese MOS-Struktur stellt einen Kondensator dar, der keinen Source- und Drainbereich hat. Die Ladungsverteilung wird durch die Spannung zwischen der Metallelektrode (Gate) und dem Siliziumsubstrat bestimmt. Hier soll nur der Fall des p-dotierten Siliziums diskutiert werden. Folgende 4 Fälle sind möglich:

- a.) Bei negativer Spannung am Gate sammeln sich an der Oberfläche des Si-Substrates

positiv geladene Löcher, die im p-dotierten Si-Substrat reichlich vorhanden sind. Diese mit Löchern angereicherte Schicht wird als Anreicherungsschicht bezeichnet.

- b.) Ohne angelegte Gate-Spannung tritt kein elektrisches Feld im Oxid auf. Damit ist das Si-Substrat zur Außenwelt elektrisch neutral. Dieser Fall wird als Flachbandfall bezeichnet.
- c.) Bei einer relativ kleinen positiven Spannung am Gate drückt das elektrische Feld die beweglichen Löcher im Si-Substrat von der Oberfläche weg. Somit bildet sich unter dem Oxid eine Raumladungszone. In dieser Raumladungszone sind die negativ geladenen Akzeptoren feststehend. Mit zunehmender positiver Gate-Spannung nimmt die Dicke der Raumladungszone zu.
- d.) Ab einer bestimmten Größe der positiven Gatespannung bildet sich an der Oberfläche des Si-Substrates unter der Oxidschicht durch das starke elektrische Feld eine Schicht aus beweglichen Elektronen. Bei weiterem Anstieg der Gatespannung dehnt sich die Raumladungszone nicht mehr aus, sondern die Dichte der beweglichen Elektronen in dieser Oberflächenschicht nimmt zu. Diese Schicht wird als Inversionsschicht bezeichnet.

3.4.2 Funktion und Kennlinien eines MOSFET

Diese Erläuterungen gelten auch für die Ladungsträgerverteilung im MOSFET. Je nach Dotierung ist ein unterschiedliches Verhalten des MOSFET möglich. Bild 3.65 zeigt oben die schematische Darstellung eines MOSFETs vom Verarmungstyp mit einem n- und p-Kanal. In diesem MOSFET ist bereits bei Gatespannungen von $U_{GS} = 0$ V ein leitfähiger Kanal vorhanden. Deshalb wird er auch als selbstleitender MOSFET bezeichnet. Bild 3.65 zeigt unten die schematische Darstellung eines MOSFETs vom Anreicherungstyp mit n- und p-Kanal. Bei einem Anreicherungs-MOSFET müssen erst durch eine Gatespannung bewegliche Ladungen, also Löcher oder Elektronen, im Kanal angesammelt werden, damit ein Strom zwischen Drain und Source fließen kann. Im Weiteren soll der Stromfluss in Abhängigkeit von der Gatespannung und Drain-Source-Spannung diskutiert werden. Ein Beispiel eines n-Kanal MOSFETs ist in Bild 3.66 gezeigt. Die Gatespannung ist bereits so groß, dass der Kanal im Inversionsbetrieb arbeitet. Bei einer kleinen Drain-Source-Spannung U_{DS} , fließt ein Strom im Kanal I_D , der proportional der angelegten Spannung U_{DS} ist. In diesem Bereich arbeitet der MOSFET wie ein linearer ohm'scher Widerstand. Ab einer bestimmten Spannung $U_{DS} = U_{GS} - U_{th}$ wird der Kanal abgeschnürt, d.h. die Konzentration der Ladungsträger nimmt ab (siehe Bild 3.66b). Eine weitere Zunahme von U_{DS} führt nicht zur Erhöhung von I_D , d.h. der MOSFET wird in der Sättigung betrieben. Ein geringfügiger Einfluss der Spannung U_{DS} bleibt vorhanden und wird als Kanallängenmodulation bezeichnet. Dieses Verhalten ist aus den Eingangs- und Ausgangskennlinien eines MOSFETs ersichtlich.

Die Eingangskennlinie und das Ausgangskennlinienfeld für einen selbstleitenden n-Kanal MOSFET (Verarmungstyp) sind in Bild 3.67 dargestellt. Sie sind vergleichbar

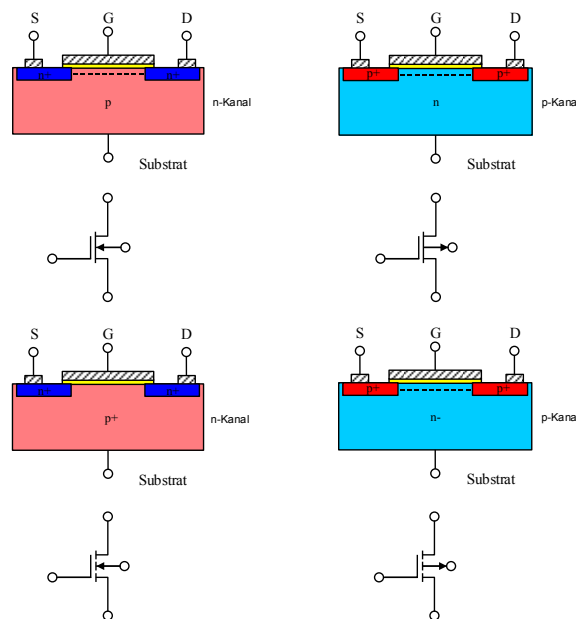


Bild 3.65: Schematische Darstellung von MOSFETs vom Verarmungstyp (oben) und Anreicherungstyp (unten) mit n- und p-Kanal.

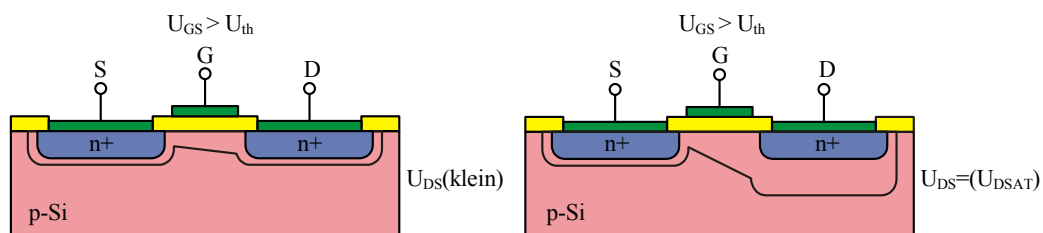


Bild 3.66: Schnitt durch einen MOSFET im linearen Betrieb (linkes Bild) und zu Beginn der Abschnürung des Kanals (rechtes Bild).

mit denen des JFET. Bei $U_{GS} = 0 \text{ V}$ fließt der Drainstrom I_{D0} . Im Ausgangskennlinienfeld sind der lineare Bereich und der Sättigungsbereich gezeigt. Im Abschnür- oder Sättigungsbereich ist der Strom I_D nur von U_{GS} abhängig. Die bereits erwähnte Kanallängenmodulation führt dazu, dass der Drainstrom leicht mit wachsender Spannung U_{DS} zunimmt. Dieser Effekt ist mit dem Early-Effekt beim bipolaren Transistor vergleichbar und wird deshalb auch mit der Early-Spannung beschrieben. Die Eingangskennlinie und das Ausgangskennlinienfeld eines selbstsperrenden n-Kanal MOSFET (Anreicherungstyp) sind in Bild 3.67 dargestellt. Bei $U_{GS} = 0 \text{ V}$ fließt hier noch kein Strom. Erst eine positive Spannung am Gate $U_{GS} > U_{th}$ führt zu einem Strom I_D . Berücksichtigen wir die Verteilung der einzelnen Ladungen und integriert man das Potential im Kanal über die Kanallänge erhält man für den Drainstrom eine Gleichung, die bereits beim JFET eingesetzt wurde. Die Gleichung der Eingangskennlinie ist auch die des JFET:

$$I_D = I_{D0} \cdot \left(1 - \frac{U_{GS}}{U_{th}}\right)^2 \quad (3.114)$$

Für $U_{GS} = 0$ V wird damit auch beim selbstleitenden MOSFET der Drainstrom zu $I_D = I_{D0}$ und für $U_{GS} = U_{th}$ wird $I_D = 0$. Im linearen Bereich und im Sättigungsbereich gelten auch hier:

$$I_D = \beta \cdot \left[(U_{GS} - U_{th}) \cdot U_{DS} - \frac{(U_{DS})^2}{2} \right]$$

und

$$I_D = \frac{1}{2} \cdot \beta \cdot (U_{GS} - U_{th})^2$$
(3.115)

Zusammenfassend erhält man unter Berücksichtigung des Early-Effekts folgende Großsignalgleichungen für Feldeffekttransistoren:

$$I_D = \begin{cases} 0 & U_{GS} \leq U_{th} \\ \beta \cdot \left((U_{GS} - U_{th}) \cdot U_{DS} - \frac{U_{DS}^2}{2} \right) \cdot \left(1 + \frac{U_{DS}}{U_A} \right) & \text{linearer Bereich} \\ \frac{\beta}{2} \cdot (U_{GS} - U_{th})^2 \cdot \left(1 + \frac{U_{DS}}{U_A} \right) & \text{Sättigungsbereich} \end{cases} \quad (3.116)$$

Der Steilkoeffizient β ist ein Maß für die Steigung der Übertragungskennlinie und wird durch die geometrischen und elektronischen Eigenschaften des Kanals bestimmt. So gilt z.B. für einen n-Kanal MOSFET:

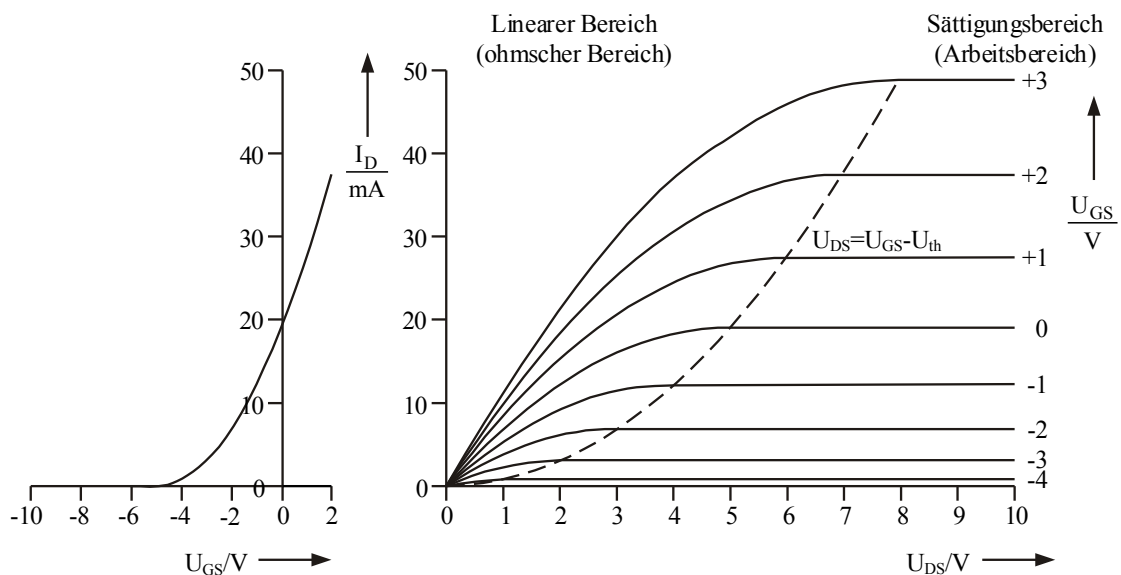


Bild 3.67: Eingangskennlinie und Ausgangskennlinienfeld eines n-Kanal MOSFET vom Verarmungstyp.

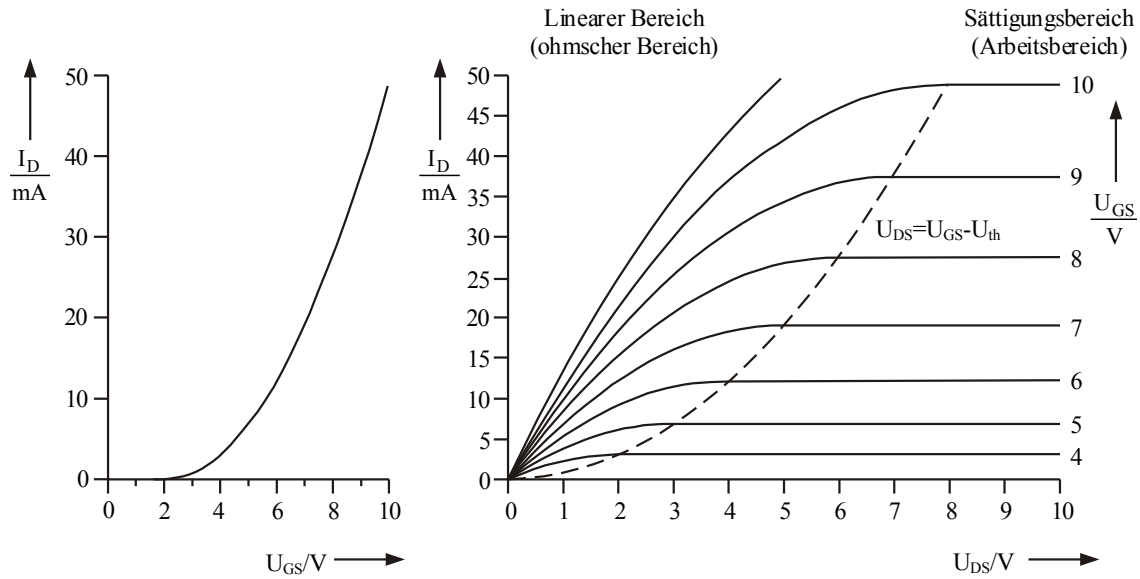


Bild 3.68: Eingangskennlinie und Ausgangskennlinienfeld eines n-Kanal MOSFET vom Anreicherungstyp.

$$\beta = \beta'_n \cdot \frac{w}{\ell} = \mu_n \cdot C'_{ox} \cdot \frac{w}{\ell} \quad (3.117)$$

Dabei ist C'_{ox} der Kapazitätsbelag des Gateoxids. Damit gilt $C_{ox} = C'_{ox} \cdot w \cdot \ell$. Der nun „physikalisch“ beschriebene Steilheitskoeffizient hat die Einheit [A/V]. Beim n-Kanal MOSFET beträgt er $20 - 60 \mu\text{A}/\text{V}^2$ und bei p-Kanal MOSFET ist er bei gleichen Abmessungen aufgrund der geringeren Beweglichkeit der Löcher 2 – 3 mal kleiner. Die Definition des Steilheitskoeffizienten β nach Gl. 3.117 muss beim MOSFET vom Anreicherungstyp immer verwendet werden, da bei $U_{GS} = 0 \text{ V}$ der Drainstrom immer $I_D = 0$ ist und damit auch $\beta = 0$ wäre. Beim selbstleitenden MOSFET gilt aber auch die in Gl. 3.97 gezeigte Form:

$$\beta = 2 \cdot \frac{I_{D0}}{U_{th}^2} \quad (3.118)$$

Der Unterschied zwischen MOSFET und JFET ist im Ausgangskennlinienfeld erkennbar. Da wir beim JFET mit einem gesperrten pn-Übergang Gatekanal arbeiten, darf die Eingangsspannung nicht positiv werden. Beim MOSFET ist der Eingang aber vom Kanal durch ein isolierendes Oxid getrennt. Deshalb darf die Eingangsspannung auch positive Werte annehmen, so dass dieser Transistortyp sich sehr gut dazu eignet, Spannungen um 0 V zu verstärken. Somit können Schaltungen für Eingangsspannungen aufgebaut werden, die sowohl einen Gleich- wie auch einen Wechselanteil besitzen, solange der Gleichanteil klein genug bleibt. Im Bereich der Leistungstransistoren findet der selbstsperrende MOSFET-Typ als sogenannter HEX-FET in einem breiten Anwendungsspektrum Verwendung.

Vorteile wie:

- spannungsgesteuerter Eingang,
- hohe Ströme,
- gute Eigenschaften bei Impulsbelastung und
- sehr kleiner Widerstand $r_{DS,on}$ (mΩ)

bieten viele Einsatzmöglichkeiten an.

Eine Zusammenfassung aller Typen von FET und ihrer Kennlinien ist im Bild 3.79 enthalten. In der modernen Digitaltechnik ist der selbstsperrende MOSFET das wichtigste Bauelement überhaupt. Durch eine immer weiter fortschreitende Miniaturisierung erreicht man heute bei höchstintegrierten Schaltungen bereits Taktfrequenzen im einstelligen GHz-Bereich.

3.4.3 Die Grundsaltungen

3.4.3.1 Die Sourceschaltung

Bild 3.69a zeigt die Sourceschaltung und Bild 3.69b zeigt das entsprechende Ersatzschaltbild. Bei Vernachlässigung des Early-Effekts gilt im Arbeitsbereich des Transistors:

$$I_D = \frac{\beta}{2} \cdot (U_{GS} - U_{th})^2 \quad (3.119)$$

Für die Ausgangsspannung erhält man mit $U_g = U_e = U_{GS}$:

$$U_a = U_{DS} \stackrel{I_a=0}{=} U_b - I_D \cdot R_D = U_b - \frac{R_D \cdot \beta}{2} \cdot (U_e - U_{th})^2 \quad (3.120)$$

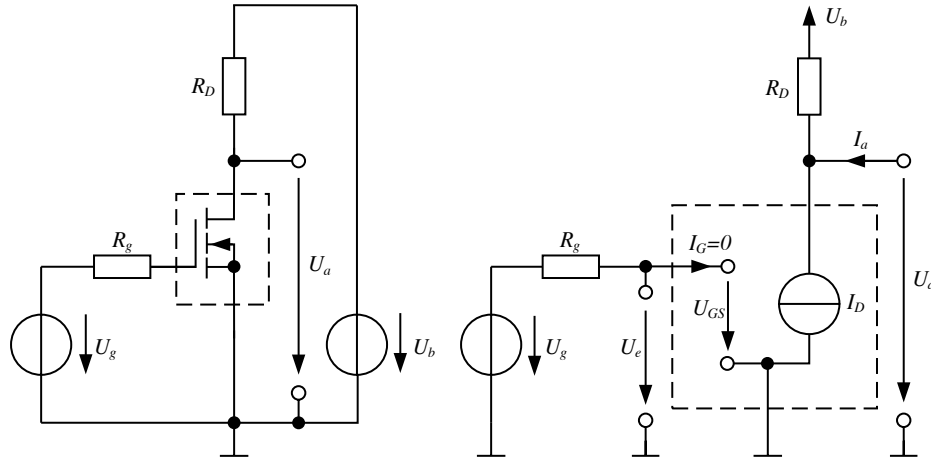


Bild 3.69: Sourceschaltung (a) und Ersatzschaltbild (b).

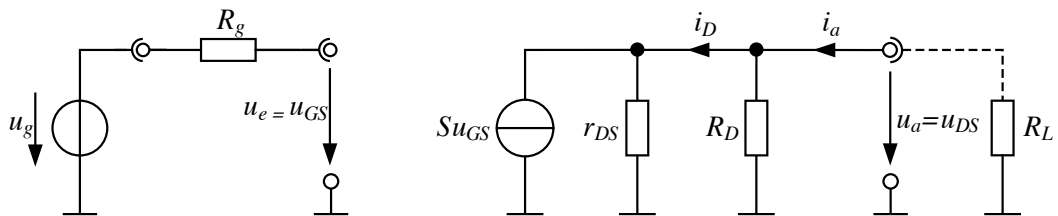


Bild 3.70: Kleinsignal-Ersatzschaltbild der Sourceschaltung.

Der Innenwiderstand R_g der Quelle hat bei MOSFETs wegen $I_G = 0$ keinen Einfluss auf die Kennlinie. Er wirkt sich nur auf das dynamische Verhalten aus. Bei Sperrschicht-FETs treten dagegen Gate-Leckströme im pA- bzw. nA-Bereich auf, die bei sehr hohen Innenwiderständen der Spannungsquelle einen nicht mehr vernachlässigbaren Spannungsabfall zur Folge haben. Deshalb setzt man bei Quellen mit $R_g > 10 \text{ M}\Omega$ bevorzugt MOSFETs ein.

Kleinsignalverhalten der Sourceschaltung

Das Verhalten bei Aussteuerung um einen Arbeitspunkt A wird als Kleinsignalverhalten bezeichnet. Bild 3.70 zeigt das Kleinsignal-Ersatzschaltbild der Sourceschaltung, das man durch Einsetzen des Kleinsignal-Ersatzschaltbildes des FETs erhält. Ohne Lastwiderstand R_L folgt aus Bild 3.70 für die Sourceschaltung:

$$A = \left. \frac{u_a}{u_e} \right|_{i_a=0} = -S \cdot (R_D \parallel r_{DS}) \stackrel{r_{DS} \gg R_D}{\approx} -S \cdot R_D \quad (3.121)$$

$$r_e = \frac{u_e}{i_e} = \infty \quad (3.122)$$

$$r_a = \frac{u_a}{i_a} = R_D \parallel r_{DS} \stackrel{r_{DS} \gg R_D}{\approx} R_D \quad (3.123)$$

Die Größen A , r_e und r_a beschreiben die Sourceschaltung vollständig. Mit Hilfe von Bild 3.70 kann man die Kleinsignal-Betriebsverstärkung berechnen:

$$A_B = \frac{u_a}{u_e} = \frac{r_e}{r_e + R_g} \cdot A \cdot \frac{R_L}{R_L + r_a} \stackrel{r_e \rightarrow \infty}{=} A \cdot \frac{R_L}{R_L + r_a} \quad (3.124)$$

Sie setzt sich aus der Verstärkung A der Schaltung und dem Spannungsteiler-Faktor am Ausgang zusammen. Die maximale Verstärkung hängt vom Arbeitspunkt ab. Sie nimmt mit zunehmendem Strom bzw. zunehmender Spannung $U_{GS} - U_{th}$ ab. Will man eine hohe maximale Verstärkung erreichen, muss man einen MOSFET mit möglichst großem Steilheitskoeffizienten β und möglichst kleinem Strom $I_{D,A}$ betreiben.

Sourceschaltung mit Stromgegenkopplung

Die Nichtlinearität und die Temperaturabhängigkeit der Sourceschaltung können durch eine Stromgegenkopplung verringert werden. Dazu wird ein Sourcewiderstand R_S eingefügt (vgl. Bild 3.71a). Die Übertragungskennlinie und das Kleinsignalverhalten hängen in diesem Fall von der Beschaltung des Substratanschlusses ab. Er ist bei Einzel-MOSFETs mit der Source und in integrierten Schaltungen mit der negativsten Versorgungsspannung (hier: Masse) verbunden. In Bild 3.71a ist deshalb ein Umschalter für den Substratanschluss enthalten. Bild 3.71b zeigt das Ersatzschaltbild. Für den Abschnürbereich erhält man mit $I_a = 0$:

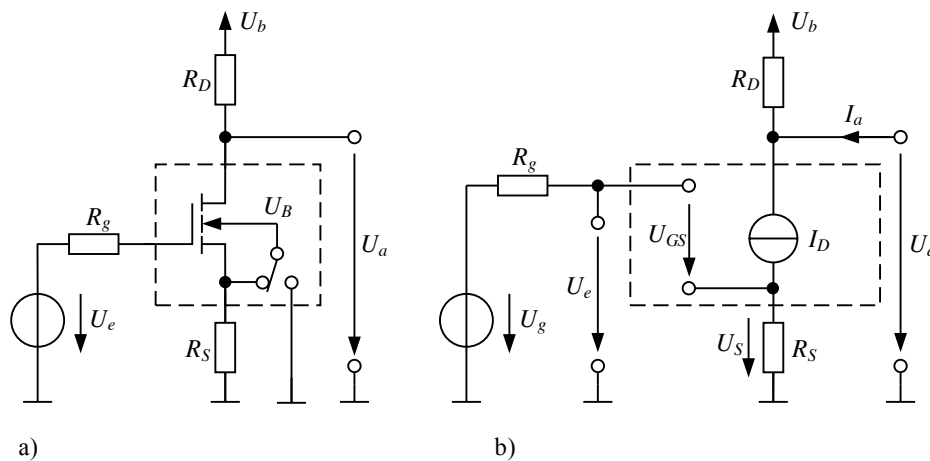


Bild 3.71: Sourceschaltung mit Stromgegenkopplung (a) und Großsignal-Ersatzschaltbild (b).

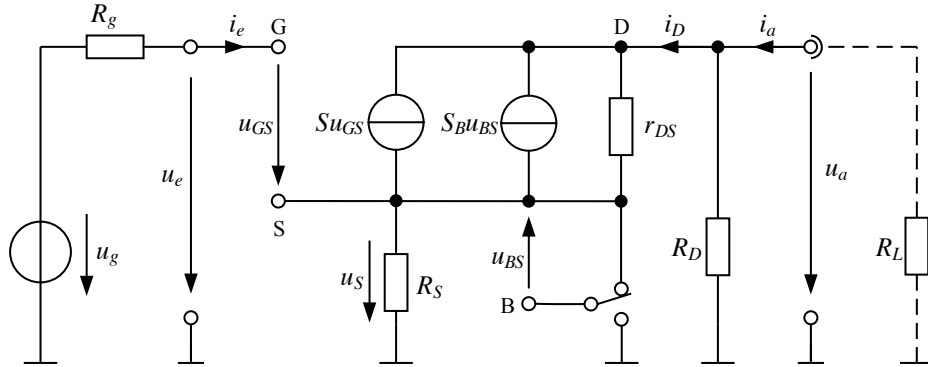


Bild 3.72: Kleinsignal-Ersatzschaltbild der Sourceschaltung mit Stromgegenkopplung.

$$U_a = U_b - I_D \cdot R_D = u_b - \frac{R_D \cdot \beta}{2} \cdot (U_{GS} - U_{th})^2 \quad (3.125)$$

$$U_e = U_{GS} + U_S = U_{GS} + I_D \cdot R_S \quad (3.126)$$

Das Kleinsignalverhalten berechnet sich anhand des Kleinsignal-Ersatzschaltbildes (siehe Bild 3.72). Aus der folgenden Knotengleichung

$$S \cdot u_{GS} + S_B \cdot u_{BS} + \frac{u_{DS}}{r_{DS}} + \frac{u_a}{R_D} = 0 \quad (3.127)$$

erhält man mit $u_{GS} = u_e - u_S$ und $u_{DS} = u_a - u_S$ die Verstärkung:

$$A = \frac{u_a}{u_e} \bigg|_{i_a=0} = - \frac{S \cdot R_D}{1 + \frac{R_D}{r_{DS}} + \left(S + S_B + \frac{1}{r_{DS}} \right) \cdot R_S} \quad (3.128)$$

$$A \stackrel{r_{DS} \gg R_D, 1/S}{\approx} = - \frac{S \cdot R_D}{1 + (S + S_B) \cdot R_S} \quad (3.129)$$

$$A \stackrel{u_{BS}=0}{=} - \frac{S \cdot R_D}{1 + S \cdot R_S} \stackrel{SR_S \gg 1}{\approx} - \frac{R_D}{R_S} \quad (3.130)$$

Bei Einzel-MOSFET, d.h. ohne Substrateffekt und starker Gegenkopplung, hängt die Verstärkung nur noch von R_D und R_S ab. Allerdings kann man aufgrund der geringen Maximalverstärkung eines MOSFETs im Allgemeinen keine starke Gegenkopplung vornehmen, weil sonst die Verstärkung zu klein wird. Bei Betrieb mit einem Lastwiderstand R_L kann man die zugehörige Verstärkung berechnen, indem man die Parallelschaltung von R_D und R_L einsetzt. Für den Eingangswiderstand gilt $r_e = \infty$ und für den Ausgangswiderstand ergibt sich:

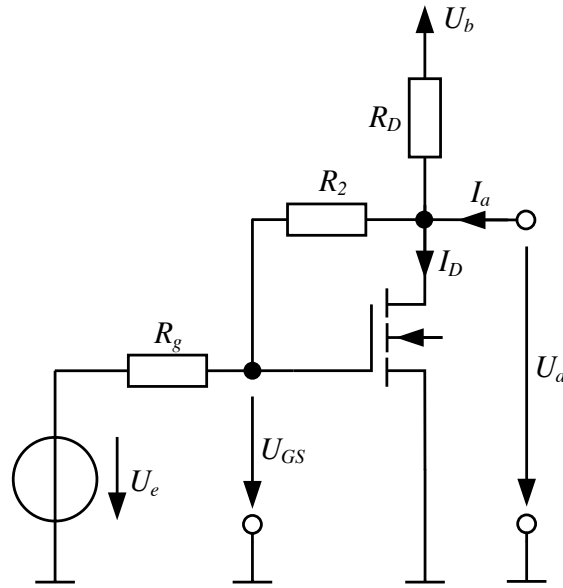


Bild 3.73: Sourceschaltung mit Spannungsgegenkopplung.

$$r_a = R_D \parallel r_{DS} \cdot \left(1 + \left(S + S_B + \frac{1}{r_{DS}} \right) \cdot R_S \right) \stackrel{r_{DS} \gg R_D}{\approx} R_D \quad (3.131)$$

Mit $r_{DS} \gg R_D$ und $r_{DS} \gg \frac{1}{S}$ sowie ohne Lastwiderstand erhält man für die Sourceschaltung mit Stromgegenkopplung folgende Gleichungen:

$$A = \frac{u_a}{u_e} \bigg|_{i_a=0} \approx - \frac{S \cdot R_D}{1 + (S + S_B) \cdot R_S} \stackrel{u_{BS}=0}{=} - \frac{S \cdot R_D}{1 + S \cdot R_S} \quad (3.132)$$

$$r_e = \infty \quad (3.133)$$

$$r_a = \frac{u_a}{i_a} \approx R_D \quad (3.134)$$

Sourceschaltung mit Spannungsgegenkopplung

Bei der Sourceschaltung mit Spannungsgegenkopplung nach Bild 3.73 wird ein Teil der Ausgangsspannung über die Widerstände R_1 und R_2 auf das Gate des FETs zurückgeführt. Aus der Knotengleichung ergibt sich

$$\frac{U_b - U_a}{R_D} + I_a = I_D + \frac{U_a - U_{GS}}{R_2} \quad (3.135)$$

$$\frac{U_{GS} - U_e}{R_1} = \frac{U_a - U_{GS}}{R_2} \quad (3.136)$$

Für den Betrieb ohne Last, d.h. es gilt $I_a = 0$

$$U_a = \frac{U_b \cdot R_2 - I_D \cdot R_D \cdot R_2 + U_{GS} \cdot R_D}{R_2 + R_D} \stackrel{R_2 \gg R_D}{\approx} U_b - I_D \cdot R_D \quad (3.137)$$

$$U_e = \frac{U_{GS} \cdot (R_1 + R_2) - U_a \cdot R_1}{R_2} \quad (3.138)$$

Das Kleinsignalverhalten errechnet sich anhand von Bild 3.71. Aus den Knotengleichungen ergibt sich

$$\frac{u_e - u_{GS}}{R_1} + \frac{u_a - u_{GS}}{R_2} = 0 \quad (3.139)$$

$$S \cdot u_{GS} + \frac{u_a - u_{GS}}{R_2} + \frac{u_a}{r_{DS}} + \frac{u_a}{R_D} = i_a \quad (3.140)$$

Die Spannungsverstärkung ist mit $R' = R_D \parallel r_{DS}$

$$A = \left. \frac{u_a}{u_e} \right|_{u_e=0} = \frac{-S \cdot R_2 + 1}{1 + S \cdot R_1 + \frac{R_1 + R_2}{R'_D}} \stackrel{\substack{r_{DS} \gg R_D \\ R_1, R_2 \gg 1/S}}{\approx} -\frac{R_2}{R_1 + \frac{R_1 + R_2}{S \cdot R_D}} \quad (3.141)$$

Für den Eingangswiderstand bei Leerlauf am Ausgang ($i_a = 0$) erhält man

$$r_{e,L} = \left. \frac{u_e}{i_e} \right|_{i_a=0} = R_1 + \frac{R_2 + R'_D}{1 + S \cdot R'_D} \stackrel{r_{DS} \gg R_D \gg 1/S}{\approx} R_1 + \frac{1}{S} \cdot \left(1 + \frac{R_2}{R_1} \right) \quad (3.142)$$

Für den Ausgangswiderstand bei Kurzschluss am Eingang erhält man

$$r_{a,K} = \left. \frac{u_a}{i_a} \right|_{u_e=0} = R'_D \left\| \frac{R_1 + R_2}{1 + S \cdot R_1} \right. \stackrel{\substack{r_{DS} \gg R_D \\ R_1 \gg 1/S}}{\approx} R_D \left\| \frac{1}{S} \cdot \left(1 + \frac{R_2}{R_1} \right) \right. \quad (3.143)$$

Mit $R_1 \rightarrow \infty$, also Leerlauf am Eingang wird der Ausgangswiderstand

$$r_{a,L} = \left. \frac{u_a}{i_a} \right|_{i_e=0} = R'_D \left\| \frac{1}{S} \right. \stackrel{r_{DS} \gg R_D \gg 1/S}{\approx} \frac{1}{S} \quad (3.144)$$

Zusammenfassend gilt für die Sourceschaltung mit Spannungsgegenkopplung:

$$A = \left. \frac{u_a}{u_e} \right|_{u_e=0} \approx - \frac{R_2}{R_1 + \frac{R_1+R_2}{S \cdot R_D}} \quad (3.145)$$

$$r_e = \left. \frac{u_e}{i_e} \right|_{i_a=0} \approx R_1 + \frac{1}{S} \cdot \left(1 + \frac{R_2}{R_D} \right) \quad (3.146)$$

$$r_a = \left. \frac{u_a}{i_a} \right|_{u_e=0} \approx R_D \parallel \frac{1}{S} \cdot \left(1 + \frac{R_2}{R_1} \right) \quad (3.147)$$

Arbeitspunkteinstellung

Der Betrieb als Kleinsignalverstärker erfordert eine stabile Einstellung des Arbeitspunktes. Der Arbeitspunkt sollte möglichst wenig von den Parametern des MOSFET abhängen, da diese temperaturabhängigen und fertigungsbedingten Streuungen unterworfen sind. Kleinsignalverstärker in Sourceschaltung mit Einzel-FETs werden aufgrund ihrer im Vergleich zur Emitterschaltung geringen Verstärkung nur in Ausnahmefällen eingesetzt. Dazu gehören Verstärker mit sehr hochohmigen Signalquellen, wie z.B. für Kondensatormikrofone. Die Arbeitspunkteinstellung bei Wechselspannungen erfolgt über Koppelkondensatoren mit der Signalquelle und der Last. Bei Spannungsverstärkern wird in der Regel die Gleichstrom-Gegenkopplung verwendet, wie es z.B. in Bild 3.74 gezeigt ist.

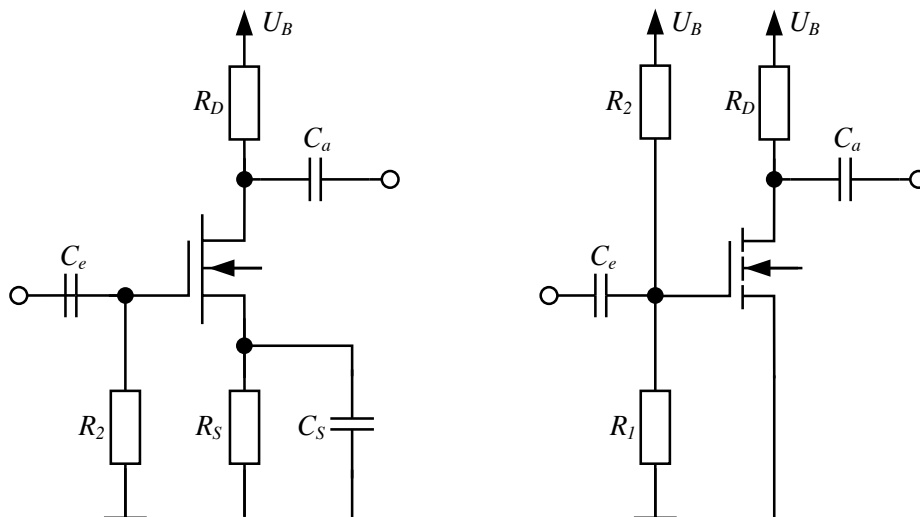


Bild 3.74: Arbeitspunkteinstellung für die Sourceschaltung für selbstleitende (a) und selbstsperrende (b) FETs.

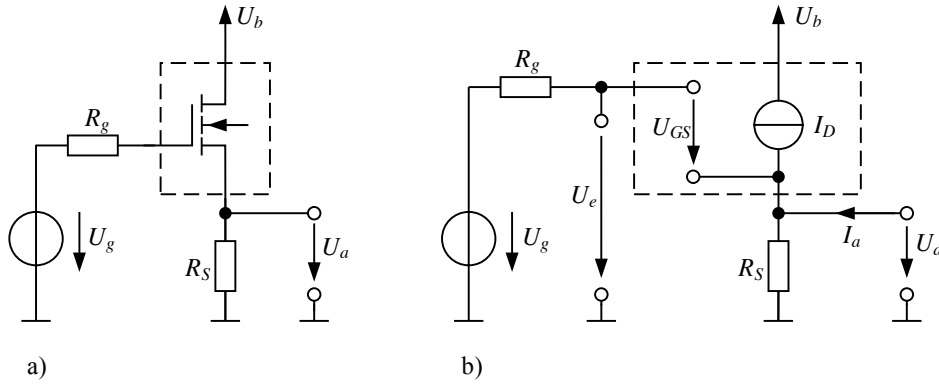


Bild 3.75: Drainschaltung (a) und Großsignal-Ersatzschaltbild (b).

3.4.3.2 Die Drainschaltung

Bild 3.75a zeigt die Drainschaltung. Das Ersatzschaltbild nach Bild 3.75b ermöglicht die Berechnung der grundlegenden Parameter der Drainschaltung für $U_g \geq U_{th}$ und $I_a = 0$ gilt:

$$U_a = I_D \cdot R_S \quad (3.148)$$

$$U_e = U_a + U_{GS} = U_a + \sqrt{\frac{2 \cdot I_D}{\beta}} + U_{th} \quad (3.149)$$

Dabei wird der Early-Effekt vernachlässigt und durch Einsetzen von Gl. 3.148 in Gl. 3.149 erhält man

$$U_e = U_a + \sqrt{\frac{2 \cdot I_D}{\beta \cdot R_S}} + U_{th} \quad (3.150)$$

Kleinsignalverhalten der Drainschaltung

Bild 3.76 zeigt im oberen Teil das Kleinsignal-Ersatzschaltbild der Drainschaltung in seiner unmittelbaren Form. Daraus erhält man durch Umzeichnen und Zusammenfassen parallel liegender Elemente das untere Kleinsignal-Ersatzschaltbild. Aus der Knotengleichung erhält man mit $u_{GS} = u_e - u_a$ die Kleinsignalverstärkung mit $R'_S = R_S \parallel r_{DS}$

$$A = \frac{u_a}{u_e} \bigg|_{i_a=0} = \frac{S \cdot R'_S}{1 + S \cdot R'_S} \stackrel{r_{DS} \gg 1/S}{\approx} \frac{S \cdot R_S}{1 + (S + S_B) \cdot R_S} \stackrel{u_{BS}=0}{=} \frac{S \cdot R_S}{1 + S \cdot R_S} \quad (3.151)$$

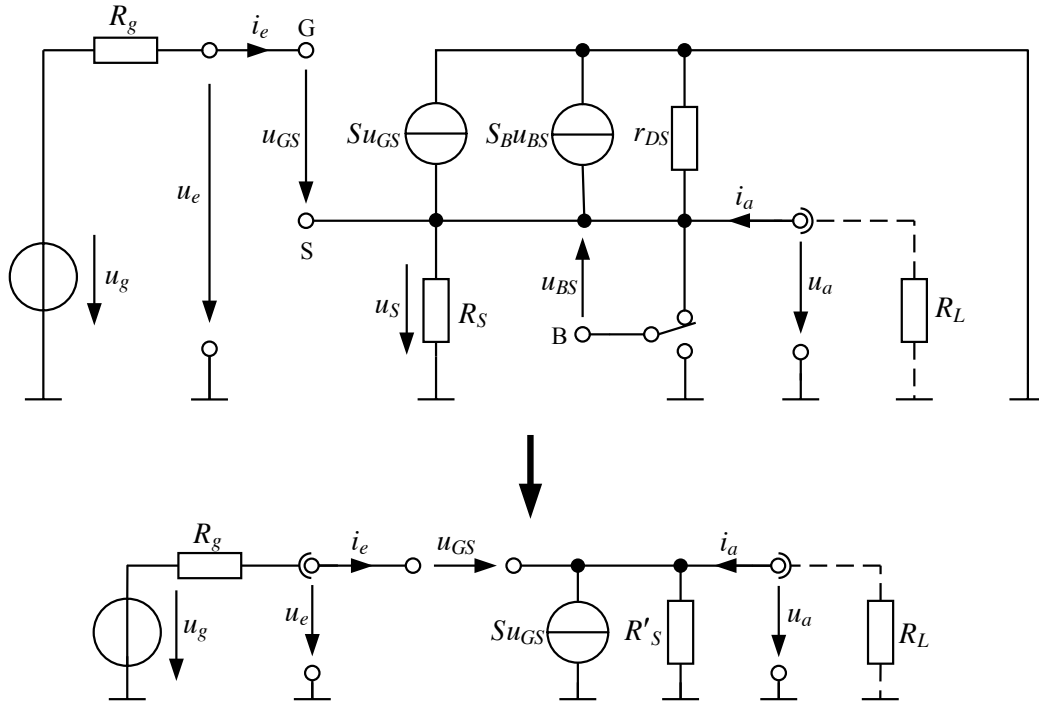


Bild 3.76: Kleinsignalersatzschaltbilder der Drainschaltung.

Für den Kleinsignal-Eingangswiderstand gilt $r_e = \infty$ und für den Kleinsignal-Ausgangswiderstand erhält man

$$r_a = \frac{u_a}{i_a} = \frac{1}{S} \left\| R'_S \right\|_{r_{DS} \gg 1/S} \approx \frac{1}{S} \left\| \frac{1}{S_B} \right\| R_S \quad (3.152)$$

Ohne Lastwiderstand und mit $r_{DS} \gg 1/S$ erhält man folgende Formeln für die Drainschaltung:

$$A = \frac{u_a}{u_e} \bigg|_{i_a=0} \approx \frac{S \cdot R_S}{1 + (S + S_B) \cdot R_S} \quad (3.153)$$

$$r_e = \frac{u_e}{i_e} \bigg|_{i_a=0} = \infty \quad (3.154)$$

$$r_a = \frac{u_a}{i_a} \approx \frac{1}{S} \left\| \frac{1}{S_B} \right\| R_S \quad (3.155)$$

Um den Einfluss des Lastwiderstandes zu berücksichtigen, muss man in Gl. 3.153 anstelle von R_S die Parallelschaltung von R_S und R_L einsetzen.

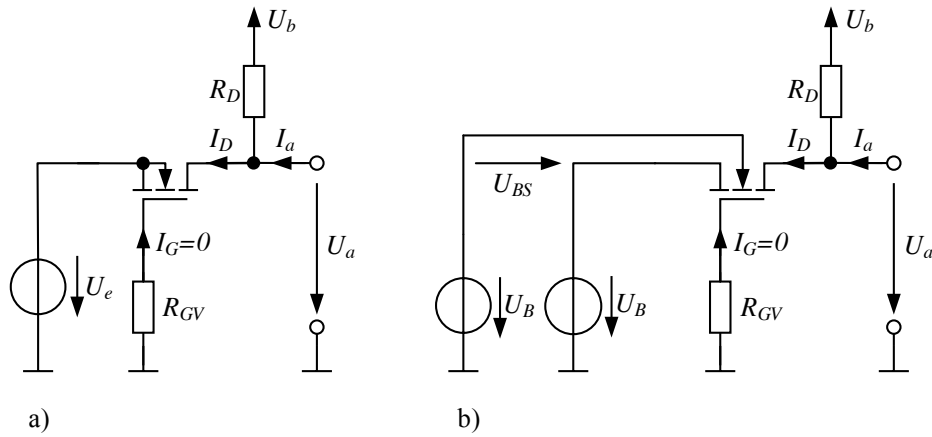


Bild 3.77: Gateschaltung a) ohne, und b) mit Anschluss einer Substratspannung.

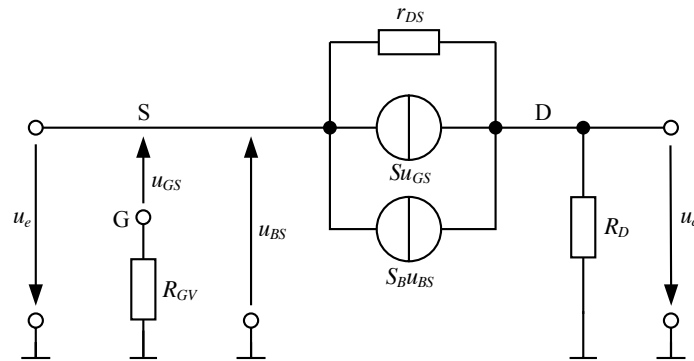


Bild 3.78: Kleinsignal-Ersatzschaltbild der Gateschaltung.

3.4.3.3 Die Gateschaltung

Bild 3.77 zeigt die Gateschaltung. Bei Vernachlässigung des Early-Effekts erhält man anhand des Ersatzschaltbildes nach Bild 3.77b folgende Formeln:

$$U_a = U_b - I_D \cdot R_D = U_b - \frac{\beta \cdot R_D}{2} \cdot (U_{GS} - U_{th})^2 \quad (3.156)$$

$$U_e = -U_{GS} - I_G \cdot R_{GV} \stackrel{I_G=0}{=} -U_{GS} \quad (3.157)$$

Bild 3.78 zeigt das Kleinsignal-Ersatzschaltbild der Gateschaltung. Der Übergang vom integrierten zum Einzel-MOSFET erfolgt mit der Einschränkung $u_{BS} = 0$, d.h. in den Gleichungen wird $S_B = 0$ gesetzt.

$$\frac{u_a}{R_D} + \frac{u_a - u_e}{r_{DS}} + S \cdot u_{GS} + S_B \cdot u_{BS} = 0 \quad (3.158)$$

Aus der Knotengleichung folgt mit $u_e = -u_{GS} = u_{BS}$

$$A = \frac{u_a}{u_e} \bigg|_{i_a=0} = \left(S + S_B + \frac{1}{r_{DS}} \right) \cdot (R_D \parallel r_{DS}) \quad (3.159)$$

$$\stackrel{r_{DS} \gg R_D, 1/S}{\approx} (S + S_B) \cdot R_D \stackrel{u_{BS}=0}{=} S \cdot R_D$$

Für den Kleinsignal-Eingangswiderstand erhält man:

$$r_e = \frac{u_e}{i_e} \bigg|_{i_a=0} = \frac{R_D + r_{DS}}{1 + (S + S_B) \cdot r_{DS}} \stackrel{r_{DS} \gg R_D, 1/S}{\approx} \frac{1}{S + S_B} \stackrel{u_{BS}=0}{=} \frac{1}{S} \quad (3.160)$$

Für den Kleinsignal-Ausgangswiderstand erhält man:

$$r_a = \frac{u_a}{i_a} = R_D \left\| \frac{(1 + (S + S_B) \cdot R_g) \cdot r_{DS} + R_g}{1 + S_B \cdot R_g} \stackrel{r_{DS} \gg R_D}{\approx} R_D \right. \quad (3.161)$$

Er hängt vom Innenwiderstand R_g des Signalgenerators ab. In der Praxis kann man die Abhängigkeit vom Innenwiderstand der Signalquelle vernachlässigen. Ohne Lastwiderstand R_L und $r_{DS} \gg R_D, 1/S$ erhält man für die Gateschaltung folgende zusammenfassende Gleichungen:

$$A = \frac{u_a}{u_e} \bigg|_{i_a=0} \approx (S + S_B) \cdot R_D \stackrel{u_{BS}=0}{=} S \cdot R_D \quad (3.162)$$

$$r_e = \frac{u_e}{i_e} \bigg|_{i_a=0} \approx \frac{1}{S + S_B} \stackrel{u_{BS}=0}{=} \frac{1}{S} \quad (3.163)$$

$$r_a = \frac{u_a}{i_a} \approx R_D \quad (3.164)$$

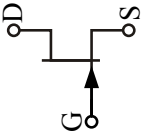
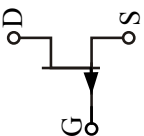
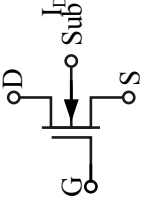
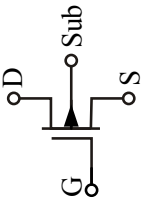
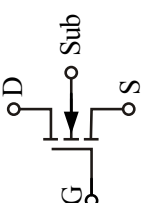
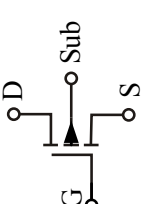
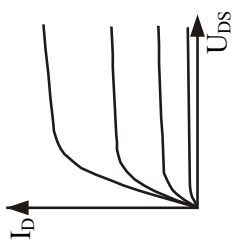
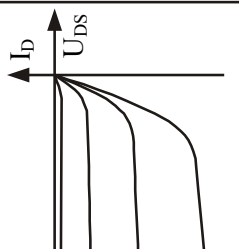
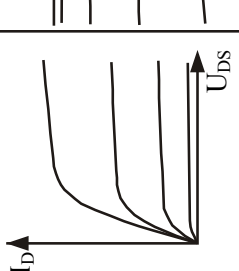
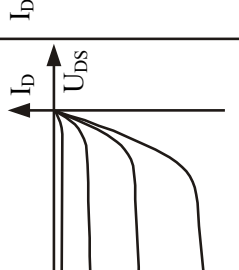
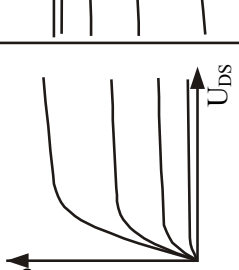
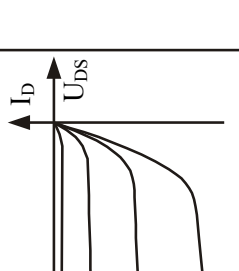
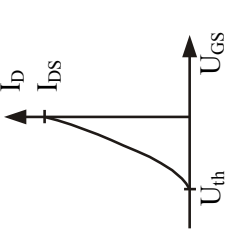
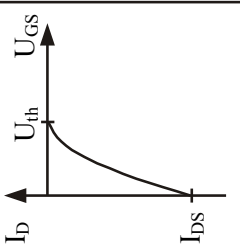
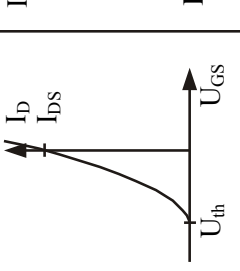
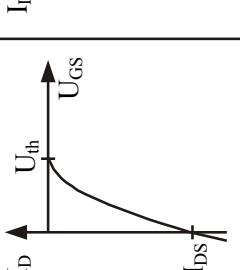
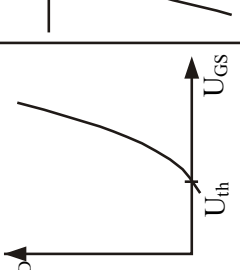
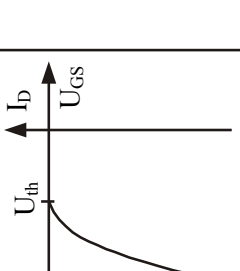
Feldeffekttransistoren					
Sperrschicht-FET (J-FET)		Isolierschicht-FET (MOSFET)			
selbstleitend (Verarmungstyp)		selbstleitend (Verarmungstyp)		selbstperrend (Anreicherungstyp)	
n-Kanal	p-Kanal	n-Kanal	p-Kanal	n-Kanal	p-Kanal
					
					
					
diskrete Verstärker	diskrete Verstärker	diskrete Hochfrequenzverstärker	diskrete Hochfrequenzverstärker	diskrete Leistungsverstärker	diskrete Leistungsverstärker
analoge ICs	analoge ICs	analoge ICs	analoge ICs	digitale ICs	digitale ICs

Bild 3.79: Übersicht über die verschiedenen Typen von Feldeffekt-Transistoren und ihre Kennlinien.

3.5 Schaltungen mit komplementären Feldeffekttransistoren (CMOS)

Lernziele

- Kennenlernen der Eigenschaften von CMOS Schaltungen
- Kennlinien von CMOS Schaltungen

3.5.1 Aufbau und Wirkungsweise

Nachdem in den vorangegangenen Kapiteln die verschiedenen Typen von Feldeffekttransistoren und deren Grundsaltungen ausführlich behandelt wurden, soll noch eine weitere Schaltungsart mit FETs betrachtet werden, die CMOS-Schaltung. Bei dieser Technik werden komplementäre Transistoren, d.h. ein n-Kanal und ein p-Kanal MOS-FET, paarweise eingesetzt. Die einfachste CMOS-Schaltung mit zwei selbstsperrenden Feldeffekttransistoren ist in Bild 3.80 dargestellt.

Zwischen der Versorgungsspannung U_b und dem Masseanschluss werden ein n-Kanal und ein p-Kanal Transistor in Reihe geschaltet. Der Sourceanschluss des p-Kanal Transistors liegt dabei direkt an U_b , die Drainanschlüsse der beiden Transistoren werden miteinander verbunden und bilden den Ausgang der Schaltung und der Sourceanschluss des n-Kanal Transistors liegt an Masse. Am Eingang der Schaltung werden die Gates der beiden Transistoren miteinander verbunden. Da Spannungen und Ströme eines p-Kanal Transistors umgekehrte Vorzeichen wie die eines n-Kanal Transistors haben, fließt also beim p-Kanal Transistor der konventionelle Strom entgegen der positiven Zählrichtung von Source nach Drain, so dass das Potential im Kanal unter dem Gate negativer ist als das Sourcepotential und daher einer negativen Gatespannung $U_{GS} < 0$ entspricht. Das Kennlinienfeld eines p-Kanal Transistors liegt bekanntlich im dritten, das des n-Kanal

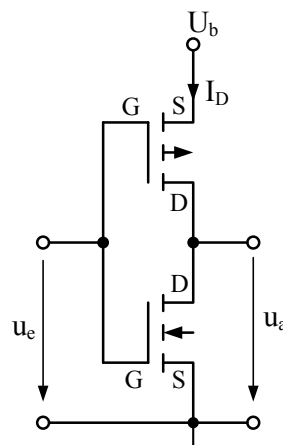


Bild 3.80: Die CMOS-Grundsaltung.

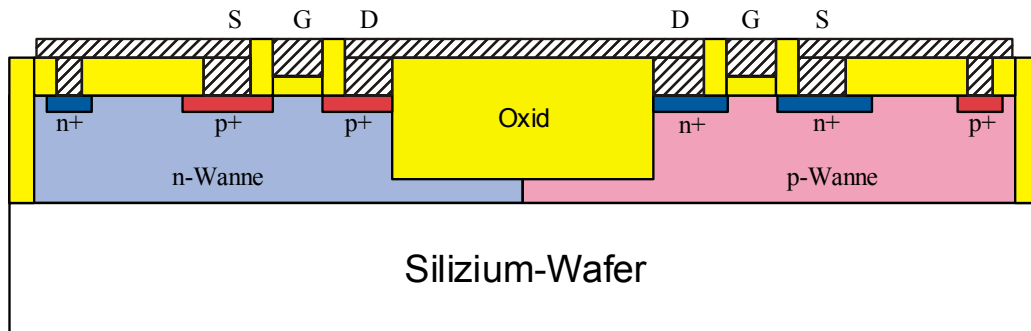


Bild 3.81: Querschnitt durch eine CMOS-Grundsaltung.

Transistoren im ersten Quadranten der I_D , U_{DS} -Ebene. Bild 3.81 zeigt einen Querschnitt eines solchen Transistorpaars. Auf einem neutralen Siliziumwafer werden die Bereiche für die beiden Transistoren definiert und mit einer entsprechenden Dotierung versehen. In den p-dotierten Bereich werden hochdotierte n-Bereiche für Source und Drain des n-Kanal Transistors und in den n-dotierten Bereich werden hochdotierte p-Bereiche für Source und Drain des p-Kanal Transistors eingebracht. Zwischen den beiden Transistoren wird an der Oberfläche ein Bereich mit einem Oxid versehen, das die Entstehung parasitärer Effekte verhindern soll. Unter dem Oxid bleibt jetzt nur noch ein kleiner Bereich übrig, in dem sich die n- und die p-Wanne berühren. Unter Berücksichtigung der an diesen Bereichen angelegten Spannungen (p-Wanne: Masse, n-Wanne: positive Versorgungsspannung) ist dieser pn-Übergang immer gesperrt, d.h. es fließt nur ein sehr kleiner Sperrstrom. Bei den neuesten integrierten CMOS-Schaltungen sind die beiden Transistoren völlig voneinander getrennt.

3.5.2 Funktion und Kennlinien

Zum Verständnis der Funktion der CMOS-Schaltung sollen die Kennlinien der beiden Transistoren zunächst einzeln betrachtet werden. Die Eingangskennlinie und das Ausgangskennlinienfeld des n-Kanal Transistors zeigt Bild 3.82 und die entsprechenden Kennlinien für den p-Kanal Transistor Bild 3.83. Die Gleichungen der Eingangskennlinie und die des Ausgangskennlinienfeldes wurden bereits im Kapitel 3.4 ausführlich beschrieben. Für den n-Kanal Transistor gilt:

$$I_D = \begin{cases} \frac{\beta}{2} \cdot (U_{GS} - U_{th})^2 & \text{Eingangskennlinie} \\ \beta \cdot \left[(U_{GS} - U_{th}) \cdot U_{DS} - \frac{(U_{DS})^2}{2} \right] & \text{linearer Bereich} \\ \frac{1}{2} \cdot \beta \cdot (U_{GS} - U_{th})^2 & \text{Sättigungsbereich} \end{cases} \quad (3.165)$$

Beim p-Kanal Transistor gelten die Formeln:

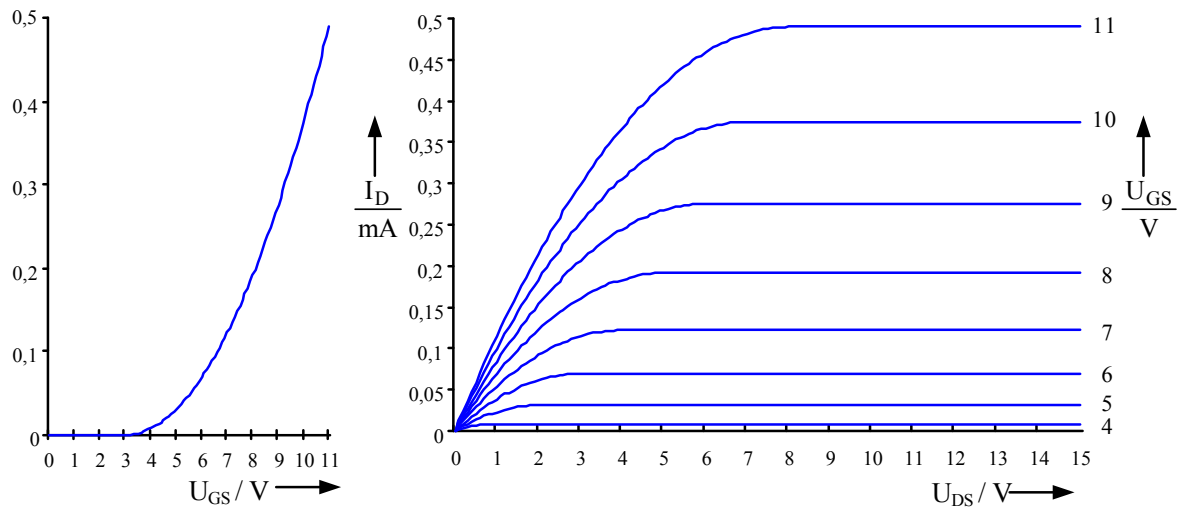


Bild 3.82: Eingangskennlinie und Ausgangskennlinienfeld des n-Kanal Transistors.

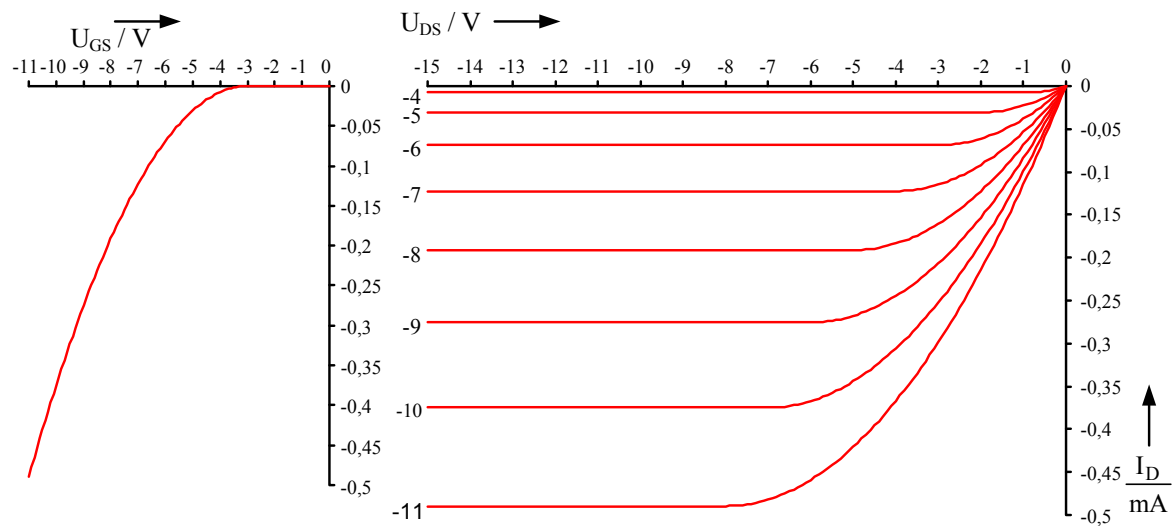


Bild 3.83: Eingangskennlinie und Ausgangskennlinienfeld des p-Kanal Transistors.

$$-I_D = \begin{cases} \frac{\beta}{2} \cdot (U_{GS} - U_{th})^2 & \text{Eingangskennlinie} \\ \beta \cdot \left[(U_{GS} - U_{th}) \cdot U_{DS} - \frac{(U_{DS})^2}{2} \right] & \text{linearer Bereich} \\ \frac{1}{2} \cdot \beta \cdot (U_{GS} - U_{th})^2 & \text{Sättigungsbereich} \end{cases} \quad (3.166)$$

Beide Transistoren haben nach den Bildern 3.82 und 3.83 betragsmäßig gleiche Kennlinien. Dies bedeutet, dass der Steilheitskoeffizient β für beide Transistoren gleich sein muss. Wie bereits erwähnt, ist die Beweglichkeit der Elektronen μ_n etwa um den Faktor 3 größer als die Beweglichkeit μ_p der Löcher. Da die Transistoren auf einer integrierten Schaltung alle gleichzeitig hergestellt werden ist zu erwarten, dass die auf die Fläche bezogene Kapazität C'_{ox} für den n- und den p-Kanal Transistor gleich groß ist. Damit kann folgender Ansatz gemacht werden:

$$\begin{aligned} \beta_n &= \mu_n \cdot C'_{ox} \cdot \frac{w_n}{\ell_n} = \beta \quad \text{und} \quad \beta_p = \mu_p \cdot C'_{ox} \cdot \frac{w_p}{\ell_p} = \beta \\ \text{mit } \mu_n &\approx 3 \cdot \mu_p \text{ und der Annahme } \ell_n = \ell_p \text{ folgt:} \\ w_p &\approx 3 \cdot w_n \end{aligned} \quad (3.167)$$

Das Ergebnis sagt aus, dass es unter den gemachten Voraussetzungen möglich ist, zwei komplementäre Feldeffekttransistoren mit betragsmäßig gleichen Eingangskennlinien und Ausgangskennlinienfeld herzustellen, wenn der p-Kanal Transistor etwa 3 mal so breit wie der n-Kanal Transistor ist.

Zur Ermittlung des bestmöglichen Arbeitspunktes der CMOS-Schaltung ist die Übergangskennlinie hilfreich. Zur Ermittlung dieser Kennlinie betrachten wir wieder die Schaltung in Bild 3.80. Der p-Kanal Transistor ersetzt hierbei den sonst üblichen Lastwiderstand. Da dieser p-Kanal Transistor aber auch mit dem Eingangssignal angesteuert wird, ist es nicht mehr auf einfache Weise möglich, einen Arbeitspunkt zu bestimmen. Als „Last“ erhalten wir jetzt das gesamte Kennlinienfeld des p-Kanal Transistors. Legen wir dieses Kennlinienfeld in das entsprechende Kennlinienfeld des n-Kanal Transistors, so erhalten wir Bild 3.84. Anstelle einer Lastgeraden erhalten wir jetzt aus den Schnittpunkten der beiden Kennlinienfelder eine nichtlineare Lastfunktion. Sie entsteht durch die unterschiedlichen Gate-Source-Spannungen an den Gates der beiden Transistoren.

Unter der Annahme, dass $U_b = 15 \text{ V}$ beträgt und die Eingangsspannung U_e von 0 V auf 15 V erhöht wird, ergeben sich für die Gate-Source-Spannungen der Transistoren die in Tabelle 3.4 aufgelisteten Spannungswerte U_{GSn} und U_{GSp} . Für die in der Tabelle 3.4 unterlegten Werte sind in Bild 3.84 die Schnittpunkte der beiden Kennlinienfelder eingetragen. Verändert man die Eingangsspannung in sehr kleinen Schritten, erhält man die in Bild 3.84 eingezeichnete „Lastkurve“. Trägt man die Ausgangsspannung U_a über der Eingangsspannung U_e auf, erhält man die Übergangskennlinie der Schaltung. Für Werte von $U_e < U_{thn}$ und $U_e > U_b - |U_{thp}|$ ist entweder der n-Kanal oder der p-Kanal

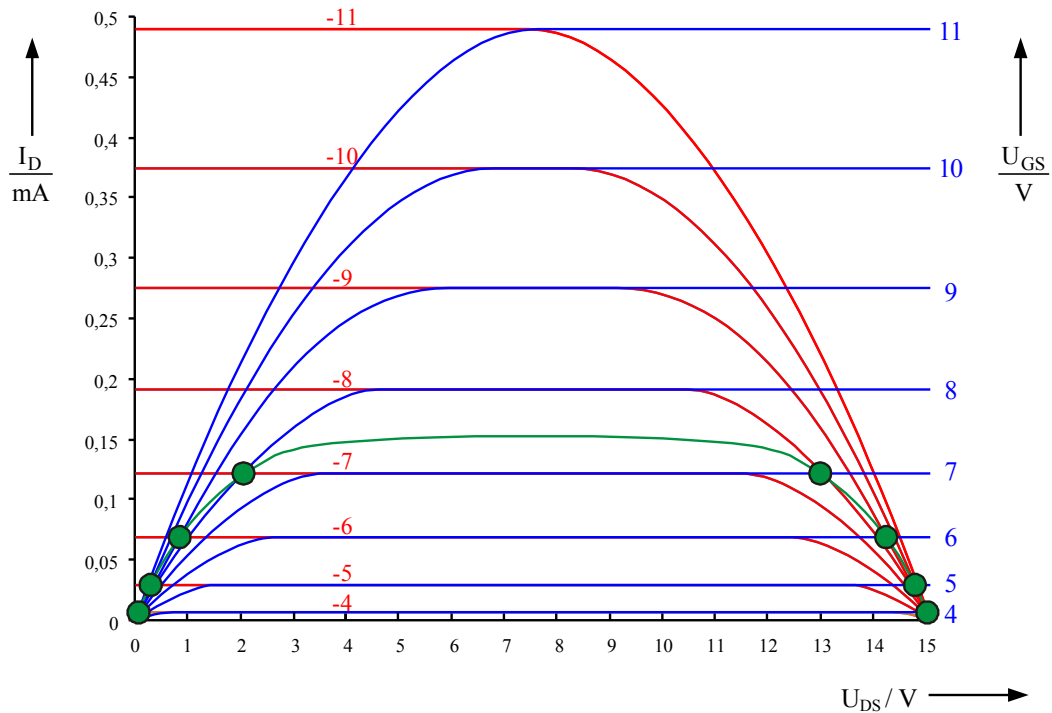


Bild 3.84: Ausgangskennlinienfeld und „Lastkurve“ der CMOS-Schaltung.

U_e [V]	U_{GSn} [V]	U_{GSp} [V]
0	0	-15
1	1	-14
2	2	-13
3	3	-12
4	4	-11
5	5	-10
6	6	-9
7	7	-8
8	8	-7
9	9	-6
10	10	-5
11	11	-4
12	12	-3
13	13	-2
14	14	-1
15	15	0

Tabelle 3.4: Zusammenhang zwischen Eingangsspannung und Gate-Source-Spannung des n-Kanal und p-Kanal Transistors einer CMOS-Grundsaltung.

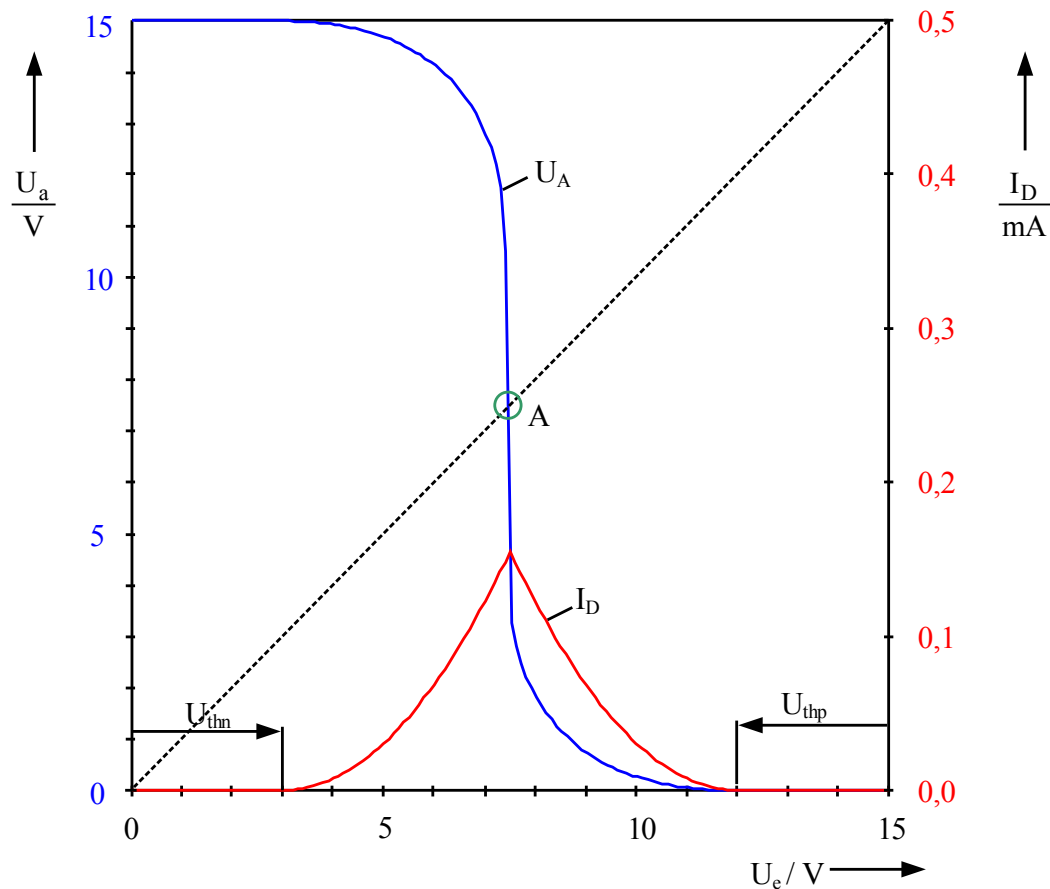


Bild 3.85: Übertragungsfunktion und Verhalten des Drainstroms einer CMOS-Schaltung.

Transistor gesperrt. Der Drainstrom ist damit $I_D = 0$. Im Bereich dazwischen nimmt der Drainstrom zunächst zu, bis beide Transistoren gleich angesteuert sind. Dies ist der Fall für $U_e = U_b/2$. Wird U_e weiter erhöht, nimmt der Strom wieder ab, bis er bei $U_e = U_b - |U_{thp}|$ den Wert 0 erreicht. Dieses Verhalten des Stromes ist ebenfalls in Bild 3.85 dargestellt. Die Steilheit der Übertragungsfunktion ist bei $U_e = U_b/2$ am größten. Deshalb ist für analoge Schaltungen in CMOS-Technik hier der ideale Arbeitspunkt. Abhängig von der Größe der Versorgungsspannung erhält man so einen mehr oder weniger großen Bereich, in dem die Schaltung nahezu linear arbeitet.

3.5.3 Die Grundsaltungen

Da bei CMOS-Schaltungen die Transistoren immer paarweise eingesetzt werden, sind Grundsaltungen wie Source-, Drain- oder Gateschaltung nicht möglich. Die einfachste Grundsaltung ist hierbei der invertierende Verstärker, wie er in Bild 3.86 in zwei Ausführungen dargestellt ist. Bild 3.86a zeigt eine Schaltung zur Verstärkung kleiner Wechsellspannungen mit nur einer positiven Spannungsquelle, während Bild 3.86b eine

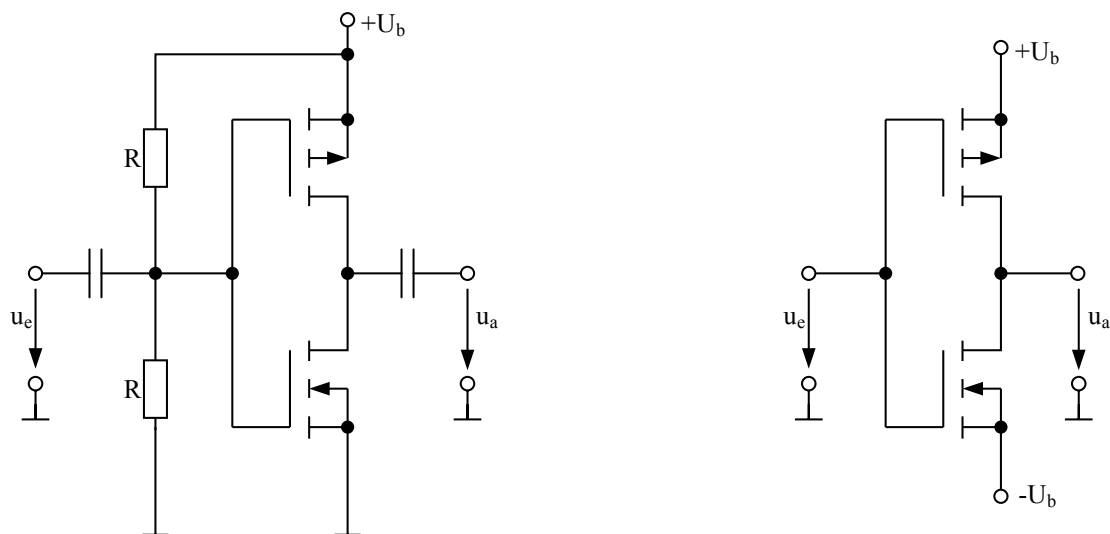


Bild 3.86: CMOS-Grundsaltungen.

Schaltung darstellt, mit der sowohl kleine Wechsel- wie auch kleine Gleichspannungen verstärkt werden können. Im Idealfall ist bei dieser Schaltung die Ausgangsspannung $u_a = 0$, wenn die Eingangsspannung $u_e = 0$ ist. Der Arbeitspunkt beider Schaltungen liegt genau bei der Hälfte der angelegten Versorgungsspannung.

4 Verstärkerschaltungen

Lernziele

- Kennenlernen der Grundlagen von Verstärkerschaltungen
- Verstehen der Funktion von Stromquellen- und Stromspiegelschaltungen
- Kennenlernen des Aufbaus und der Wirkungsweise des Differenzverstärkers
- Verstehen der Begriffe Gegenkopplung und Mitkopplung in Verstärkerschaltungen

Verstärker sind eine wesentliche Baugruppe der analogen Schaltungstechnik. In modernen Signalverarbeitungs- und Meßsystemen werden die Signale zunehmend in digitaler Form verarbeitet. Doch die Signale unserer Außenwelt sind ursächlich analoger Art und werden von Sensoren erfasst. Andere Signale werden drahtlos verbreitet, z.B. Radio, TV, Mobilfunk, WLAN. Auch diese Signale müssen zuerst von der Antenne erfasst und analog weiter verstärkt werden, bevor durch entsprechende Analog-/Digitalwandler für die digitale Weiterbearbeitung aufbereitet werden. Bild 4.1 zeigt eine Signalverarbeitungskette, in der eine physikalische Größe durch einen Sensor erfasst wird und nach analoger und digitaler Verarbeitung wieder über einen analogen Aktor auf eine physikalische Größe einwirkt. Im Signalweg in Bild 4.1 sind zwei Arten von Verstärkern eingesetzt worden. Der Verstärker nach dem Sensor arbeitet im Kleinsignalbetrieb, während der Aktor von einem Leistungsverstärker angesteuert wird.

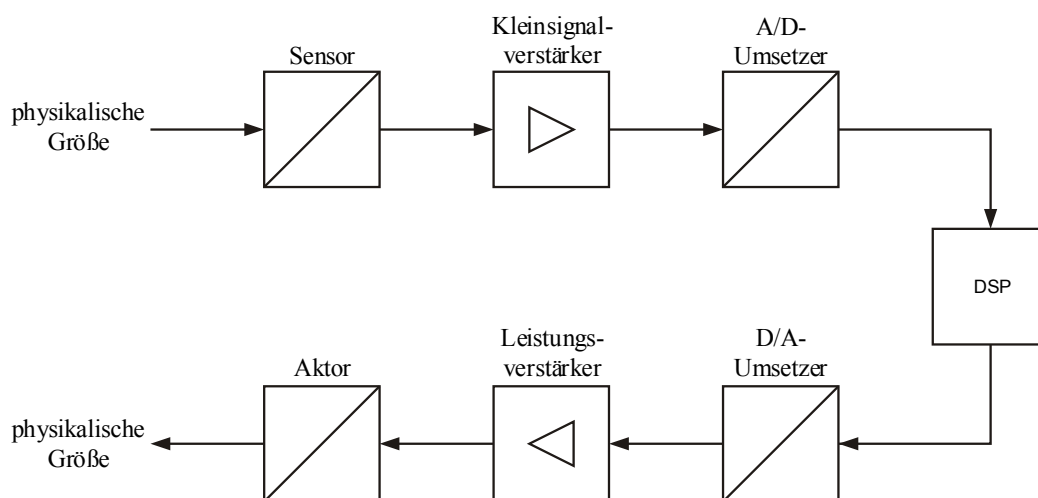


Bild 4.1: Blockschaltbild einer Signalverarbeitungskette.

Ein wichtiges Unterscheidungsmerkmal ist der Frequenzbereich, in dem der Verstärker arbeiten soll. Dabei unterscheidet man zwischen Gleichspannungs- und Wechselspannungsverstärkern. Die Wechselspannungsverstärker werden in Niederfrequenz-, Hochfrequenz- und Mikrowellenverstärker eingeteilt. Trotz ihrer Vielfalt basieren alle Verstärker auf den Transistorgrundschaltungen. Sie unterscheiden sich durch die Kopplung der Ein- und Ausgänge bzw. zwischen den einzelnen Verstärkerstufen. Eine Sonderstellung nehmen Operationsverstärker ein. Deshalb werden sie auch in Kapitel 5 extra behandelt.

4.1 Mehrstufige Verstärker

In den vorausgegangenen Kapiteln wurden die verschiedenen Grundschaltungen mit bipolaren und Feldeffekttransistoren und einfache Verstärkerschaltungen mit immer nur einem einzelnen Transistor behandelt. Oftmals ist aber die Verstärkung einer einzelnen Stufe nicht ausreichend. Dort müssen mehrstufige Verstärker eingesetzt werden. Am Beispiel eines zweistufigen Wechselspannungsverstärkers nach Bild 4.2 soll exemplarisch gezeigt werden, welche zusätzlichen Randbedingungen dabei zu beachten sind. Der Verstärker besteht aus zwei hintereinander geschalteten Emitterschaltungen. Die

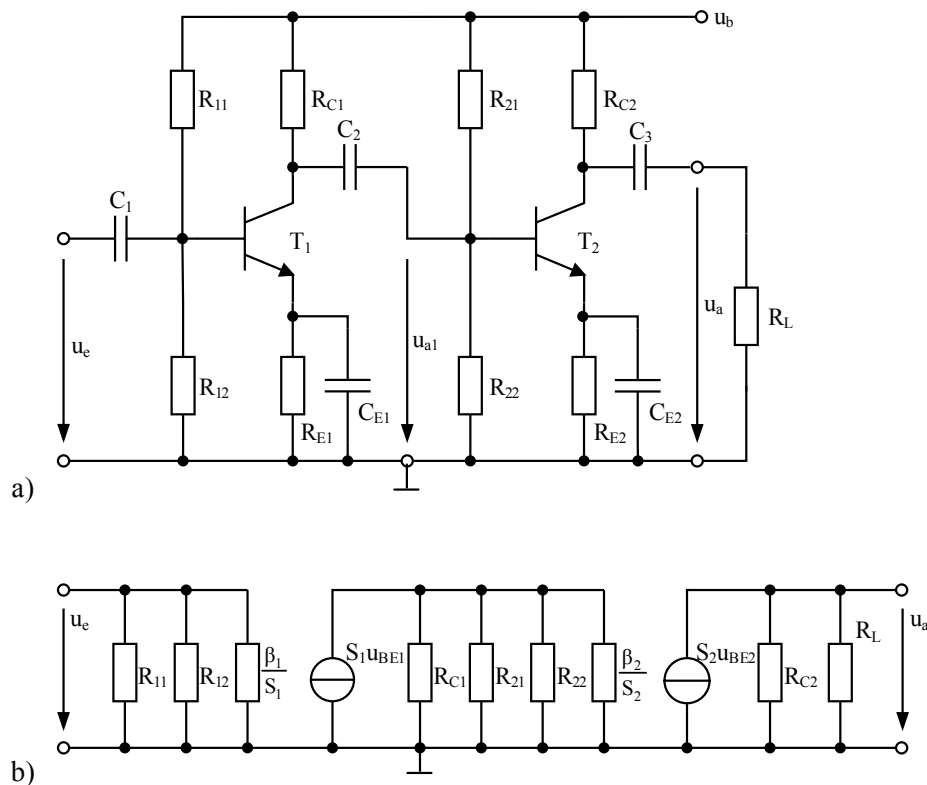


Bild 4.2: Zweistufiger Wechselspannungsverstärker. a) Gesamtschaltung, b) Kleinsignal-Ersatzschaltbild.

Verstärkung der ersten Stufe beträgt:

$$A_1 = \frac{u_{a1}}{u_e} = \frac{-i_{C1} \cdot r_{a1}}{u_{BE1}} = -S_1 \cdot r_{a1} \quad \text{mit } u_{BE1} = u_e \quad (4.1)$$

Im Widerstand r_{a1} sind alle Widerstände zusammengefasst, die für die erste Stufe als Last wirken. Dieser lässt sich anhand des Kleinsignal-Ersatzschaltbilds, Bild 4.2b, leicht ermitteln. Es ist:

$$r_{a1} = R_{C1} \parallel R_{21} \parallel R_{22} \parallel \frac{\beta_2}{S_2} \quad (4.2)$$

$$r_{a1} = \frac{R_{C1} \cdot R_{21} \cdot R_{22} \cdot \frac{\beta_2}{S_2}}{R_{21} \cdot R_{22} \cdot \frac{\beta_2}{S_2} + R_{C1} \cdot R_{22} \cdot \frac{\beta_2}{S_2} + R_{C1} \cdot R_{21} \cdot \frac{\beta_2}{S_2} + R_{C1} \cdot R_{21} \cdot R_{22}}$$

Wenn wir die bei der Emitterschaltung getroffenen Vereinbarungen, dass die Widerstände des Spannungsteilers R_{21} und R_{22} groß gegenüber dem differentiellen Basis-Emitterwiderstand β_2/S_2 des Transistors 2 im Arbeitspunkt sind, können wir die Gleichung 4.2 vereinfachen zu:

$$r_{a1} = \frac{R_{C1} \cdot \beta_2}{\beta_2 + S_2 \cdot R_{C1}} = R_{C1} \cdot \frac{1}{1 + \frac{S_2 \cdot R_{C1}}{\beta_2}} \quad (4.3)$$

Die Wechselspannungsverstärkung der ersten Stufe wird damit zu:

$$A_1 = \frac{u_{a1}}{u_e} = -S_1 \cdot R_{C1} \cdot \frac{1}{1 + \frac{S_2 \cdot R_{C1}}{\beta_2}} \quad (4.4)$$

Die Verstärkung der zweiten Stufe wird analog zu der Berechnung der ersten Stufe durchgeführt.

$$A_2 = \frac{u_a}{u_{a1}} = \frac{-i_C \cdot r_{a2}}{u_{BE2}} = -S_2 \cdot r_{a2} \quad \text{mit } u_{BE2} = u_{a1} \quad (4.5)$$

Der Lastwiderstand r_{a2} der zweiten Stufe ist nach Bild 4.2b die Parallelschaltung der beiden Widerstände R_{C2} und R_L . Damit ist:

$$r_{a2} = \frac{R_{C2} \cdot R_L}{R_{C2} + R_L} \quad (4.6)$$

Damit erhalten wir die Wechselspannungsverstärkung der zweiten Stufe:

$$A_2 = \frac{u_a}{u_{a1}} = -S_2 \cdot R_{C2} \cdot \frac{R_L}{R_{C2} + R_L} \quad (4.7)$$

Die Gesamtverstärkung beider Stufen ist das Produkt der Einzelverstärkungen A_1 und

A_2 der beiden Stufen.

$$\begin{aligned}
 A_{ges} &= \frac{u_a}{u_e} \quad \text{mit} \quad u_a = A_2 \cdot u_{a1} \quad \text{und} \quad \frac{1}{u_e} = \frac{A_1}{u_{a1}} \\
 &\quad \text{wird} \quad A_{ges} = \frac{u_a}{u_e} = A_2 \cdot u_{a1} \cdot \frac{A_1}{u_{a1}} = A_2 \cdot A_1 \\
 A_{ges} &= S_1 \cdot S_2 \cdot R_{C1} \cdot R_{C2} \cdot \frac{1}{1 + \frac{S_2 \cdot R_{C1}}{\beta_2}} \cdot \frac{R_L}{R_{C2} + R_L}
 \end{aligned} \tag{4.8}$$

Wie bei allen bisherigen Betrachtungen und Vereinfachungen müssen auch bei dieser Schaltung die Kapazitäten C_1 , C_2 , und C_3 so gewählt werden, dass ihr Blindwiderstand $1/\omega C$ für die Frequenz der zu verstärkenden Wechselspannung sehr klein wird. Je niedriger aber die Frequenzen werden, umso größer müssen die Kapazitäten der Kondensatoren werden, was wiederum auch ein größeres Volumen der Bauelemente bedeutet. Bei den modernen SMD-Bauelementen kann dies sehr schnell dazu führen, dass sinnvolle wechelspannungsgekoppelte Verstärker überhaupt nicht mehr realisierbar sind. In diesen Fällen müssen dann Gleichspannungsverstärker eingesetzt werden.

4.2 Stromquellen und Stromspiegel

Eine Stromquelle liefert einen konstanten Ausgangsstrom und ein Stromspiegel stellt am Ausgang eine verstärkte oder abgeschwächte Kopie des Eingangsstromes bereit. Damit ist ein Stromspiegel eine stromgesteuerte Stromquelle. Jeder Stromspiegel kann auch als Stromquelle betrieben werden, wenn der Eingangsstrom konstant gehalten wird.

4.2.1 Stromquellen

Betrachten wir die Ausgangskennlinien eines bipolaren Transistors oder MOSFETs, dann verlaufen diese in einem weiten Bereich horizontal (vgl. Bild 4.3). Der Kollektor- oder Drainstrom hängt faktisch nicht von der Kollektor-Emitterspannung bzw. Drain-Source-Spannung ab. Somit kann man einen Transistor als Stromquelle einsetzen, wenn man seine Eingangsspannung konstant hält. Dabei empfiehlt sich die Emitter- bzw. Sourceschaltung mit Stromgegenkopplung. Der Ausgangsstrom wird durch die Gleichungen 4.9 beschrieben:

$$I_a = \begin{cases} I_C(U_{BE}, U_{CE}) \approx I_C(U_{BE}) \stackrel{U_{BE}=\text{const.}}{=} \text{const.} \\ I_D(U_{GS}, U_{DS}) \approx I_D(U_{GS}) \stackrel{U_{GS}=\text{const.}}{=} \text{const.} \end{cases} \tag{4.9}$$

Damit erhält man nach der Maschengleichung

$$U_0 = U_{BE} + U_R = U_{BE} + (I_C + I_B) \cdot R_E \stackrel{I_C \gg I_B}{\approx} U_{BE} + I_C \cdot R_E \tag{4.10}$$

Daraus folgt mit $I_C = I_a$

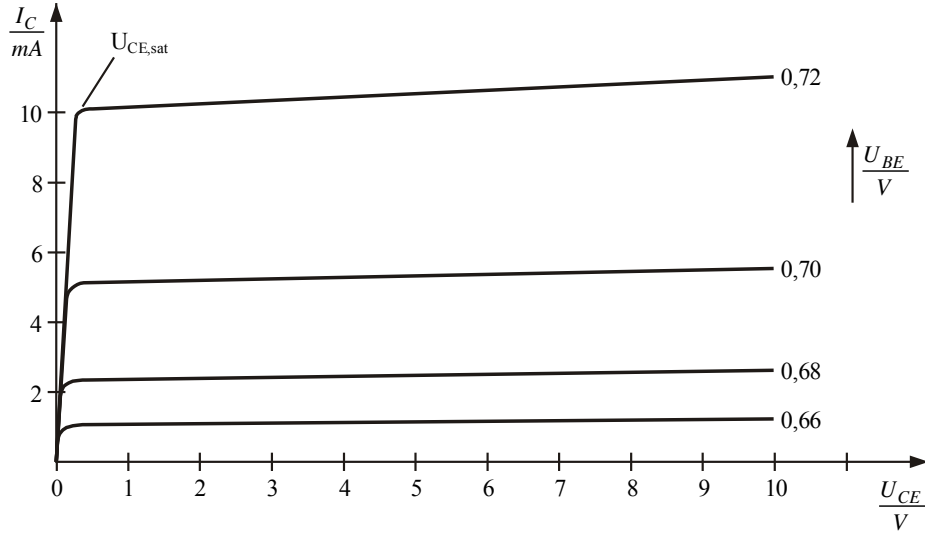


Bild 4.3: Ausgangskennlinienfeld eines bipolaren Transistors.

$$I_a \approx \frac{U_0 - U_{BE}}{R_E} \stackrel{U_{BE} \approx 0,7V}{\approx} \frac{U_0 - 0,7 V}{R_E} \quad (4.11)$$

Der Transistor in der Emitterschaltung arbeitet nur korrekt, wenn $U_{CE} > U_{CE,sat}$ ist und damit gilt:

$$U_a = U_R + U_{CE} > U_R + U_{CE,sat} = U_0 - U_{BE} + U_{CE,sat} \quad (4.12)$$

Das Ausgangskennlinienfeld ist in Bild 4.3 dargestellt. Der Innenwiderstand einer idealen Stromquelle, man kann ihn auch als Ausgangswiderstand der Stromquelle bezeichnen, ist unendlich. Der Ausgangswiderstand ist definiert als

$$r_a = \left. \frac{\partial U_a}{\partial I_a} \right|_{U_0 = \text{const.}} \quad (4.13)$$

Wie bereits diskutiert wurde, wird der reale Anstieg der Ausgangskennlinien durch den Early-Effekt verursacht. Aus der Formel für den Ausgangswiderstand der Emitterschaltung mit Stromgegenkopplung ergibt sich unter der Annahme von $R_C \rightarrow \infty$ und $R_g = 0$

$$r_a = \left. \frac{u_a}{i_a} \right|_{U_0 = \text{const.}} \stackrel{r_{CE} \gg r_{BE}}{\approx} r_{CE} \cdot \left(1 + \frac{\beta \cdot R_E}{R_E + r_{BE}} \right) \quad (4.14)$$

Für $\beta \gg 1$ und $r_{BE} = \frac{\beta}{S}$ folgt:

$$r_a \approx \begin{cases} r_{CE} \cdot (1 + S \cdot R_E) & \text{für } R_E \ll r_{BE} \\ \beta \cdot r_{CE} & \text{für } R_E \gg r_{BE} \end{cases} \quad (4.15)$$

Für eine Stromquelle mit MOSFET erhält man unter Anwendung der Sourceschaltung mit Stromgegenkopplung und $I_a = I_D$

$$U_0 = U_R + U_{GS} = I_a \cdot R_S + U_{GS} = I_a \cdot R_S + U_{th} + \sqrt{\frac{2 \cdot I_a}{\beta}} \quad (4.16)$$

Beim Einzel-MOSFET kann man I_a und U_e vorgeben und damit

$$R_S = \frac{U_0 - U_{th}}{I_a} - \sqrt{\frac{2}{\beta \cdot I_a}} \quad (4.17)$$

Den Ausgangswiderstand erhält man aus der Sourceschaltung mit Stromgegenkopplung mit:

$$r_a = \left. \frac{u_a}{i_a} \right|_{U_0=\text{const.}} \stackrel{r_{DS} \gg 1/S}{\approx} r_{DS} \cdot (1 + (S + S_B) \cdot R_S) \stackrel{S \gg S_B}{\approx} r_{DS} \cdot (1 + S \cdot R_S) \quad (4.18)$$

4.2.1.1 Stromquellen mit diskreten Transistoren

Bild 4.4 zeigt die häufigsten praktischen Stromquellen. Für $I_q \gg I_B \approx 0$ gilt

$$\left. \begin{array}{l} I_q \approx \frac{U_b}{R_1 + R_2} \\ I_q \cdot R_2 \approx I_a \cdot R_3 + U_{BE} \end{array} \right\} \Rightarrow I_a \approx \frac{1}{R_3} \cdot \left(\frac{U_b \cdot R_2}{R_1 + R_2} - U_{BE} \right) \quad \text{mit } U_{BE} \approx 0,7 \text{ V} \quad (4.19)$$

Der Ausgangsstrom ist temperaturabhängig, da die Basis-Emitterspannung von der Temperatur abhängt

$$\frac{\partial I_a}{\partial T} = -\frac{1}{R_3} \cdot \frac{\partial U_{BE}}{\partial T} \approx \frac{2 \text{ mV/K}}{R_3} \quad (4.20)$$

In der Schaltung nach Bild 4.4b wird die Temperaturabhängigkeit verringert, indem die Basis-Emitterspannung durch die Spannung an der Diode kompensiert wird. Damit gilt mit $U_D \approx U_{BE}$ und $I_q \gg I_B \approx 0$

$$\left. \begin{array}{l} I_q \approx \frac{U_b - U_D}{R_1 + R_2} \\ I_q \cdot R_2 \approx I_a \cdot R_3 \end{array} \right\} \Rightarrow I_a \approx \frac{(U_b - U_D) \cdot R_2}{(R_1 + R_2) \cdot R_3} \quad \text{mit } U_D \approx 0,7 \text{ V} \quad (4.21)$$

Somit erhält man für die Temperaturabhängigkeit des Ausgangsstromes:

$$\frac{\partial I_a}{\partial T} = -\frac{R_2}{(R_1 + R_2) \cdot R_3} \cdot \frac{\partial U_D}{\partial T} \approx \frac{2 \text{ mV/K}}{R_3} \cdot \frac{R_2}{R_1 + R_2} \approx 2 \text{ mV/K} \cdot \frac{I_a}{U_b - U_D} \quad (4.22)$$

Für die Schaltung nach Bild 4.4c ergibt sich

$$I_a \approx \frac{U_Z - U_{BE}}{R_3} \approx \frac{U_Z - 0,7 \text{ V}}{R_3} \quad (4.23)$$

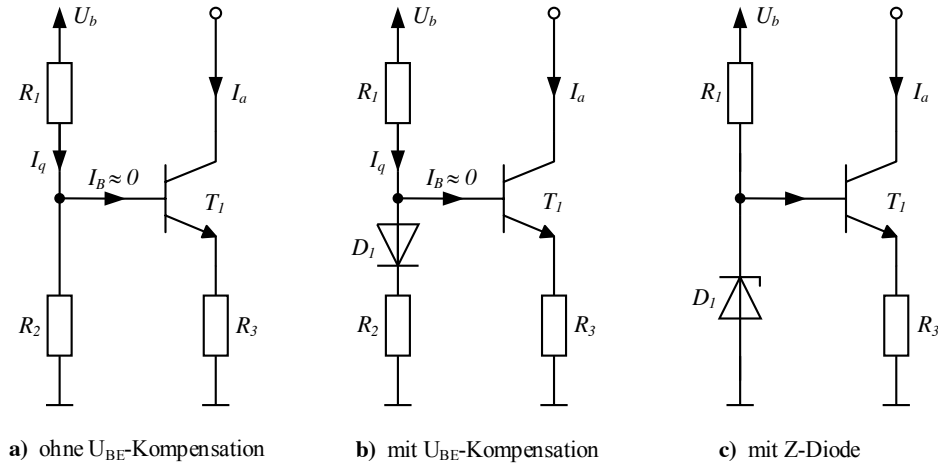


Bild 4.4: Stromquellen mit diskreten Transistoren.

Die geringste Temperaturabhängigkeit erhält man mit $U_Z \approx 5 - 6 \text{ V}$.

4.2.2 Der Stromspiegel

Bild 4.5 zeigt die Schaltung von Stromspiegeln mit bipolaren Transistoren und MOSFETs. Ein Stromspiegel besteht im einfachsten Fall aus zwei Transistoren und zwei Widerständen zur Stromgegenkopplung. Durch einen Vorwiderstand R_V kann man einen konstanten Referenzstrom einstellen und somit den Stromspiegel in eine Stromquelle umwandeln. Für den Stromspiegel mit npn Transistoren nach Bild 4.5a ergibt der Maschensatz

$$(I_{C1} + I_{B1}) \cdot R_1 + U_{BE1} = (I_{C2} + I_{B2}) \cdot R_2 + U_{BE2} \quad (4.24)$$

Nach den Formeln für die npn Transistoren nach Kapitel 3.2

$$\begin{aligned} I_{C1} &= I_{S1} \cdot e^{\frac{U_{BE1}}{U_T}} & , & \quad I_{B1} = \frac{I_{C1}}{\beta} \\ I_{C2} &= I_{S2} \cdot e^{\frac{U_{BE2}}{U_T}} \cdot \left(1 + \frac{U_{CE2}}{U_A}\right) & , & \quad I_{B2} = \frac{I_{C2}}{\beta} \end{aligned} \quad (4.25)$$

Für den ersten Transistor können wir den Early-Effekt vernachlässigen, da $U_{CE1} = U_{BE1} \gg U_A$ ist. Aus dem Knotensatz folgt

$$I_e = I_{C1} + I_{B1} + I_{B2} \quad , \quad I_a = I_{C2} \quad (4.26)$$

Für den allgemeinen Fall erhält man durch Einsetzen von Gl. 4.26 in Gl. 4.25 und bei Vernachlässigung des Early-Effekts folgende Gleichung

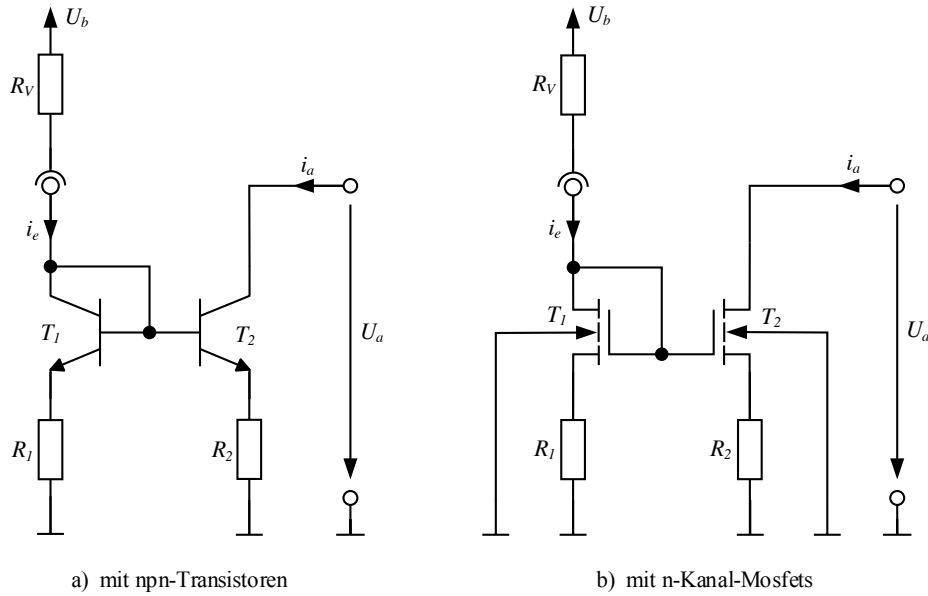


Bild 4.5: Stromspiegel mit bipolaren Transistoren und MOSFETs.

$$\left(1 + \frac{1}{B}\right) \cdot R_1 \cdot I_{C1} + U_T \cdot \ln \frac{I_{C1}}{I_{S1}} = \left(1 + \frac{1}{B}\right) \cdot R_2 \cdot I_{C2} + U_T \cdot \ln \frac{I_{C2}}{I_{S2}} \quad (4.27)$$

Diese Gleichung ist explizit nicht lösbar. Für ausreichend große Widerstände dominieren die linearen Terme in der Gleichung und man erhält:

$$R_1 \cdot I_{C1} = R_2 \cdot I_{C2} \quad (4.28)$$

Daraus folgt für den Ausgangsstrom mit Hilfe von Gl. 4.26:

$$I_a \approx \frac{R_1}{R_2 + \frac{R_1 + R_2}{B}} \cdot I_e \stackrel{B \gg 1 + R_1/R_2}{\approx} \frac{R_1}{R_2} \cdot I_e \quad (4.29)$$

Damit bestimmt nur noch das Widerstandsverhältnis das Übersetzungsverhältnis zwischen Eingangs- und Ausgangsstrom.

4.3 Der Differenzverstärker

Für viele Anwendungen werden Verstärker benötigt, die eine sehr große Bandbreite für Signale von Gleichspannung bis hin zu Frequenzen im MHz-, ja sogar bis in den GHz-Bereich besitzen müssen. Dies bedeutet, dass sowohl der Eingang der Verstärkerschaltung wie auch die einzelnen Stufen bei mehrstufigen Verstärkern, nicht mehr durch Kondensatoren entkoppelt werden können. Andere Probleme bei der Verstärkung von

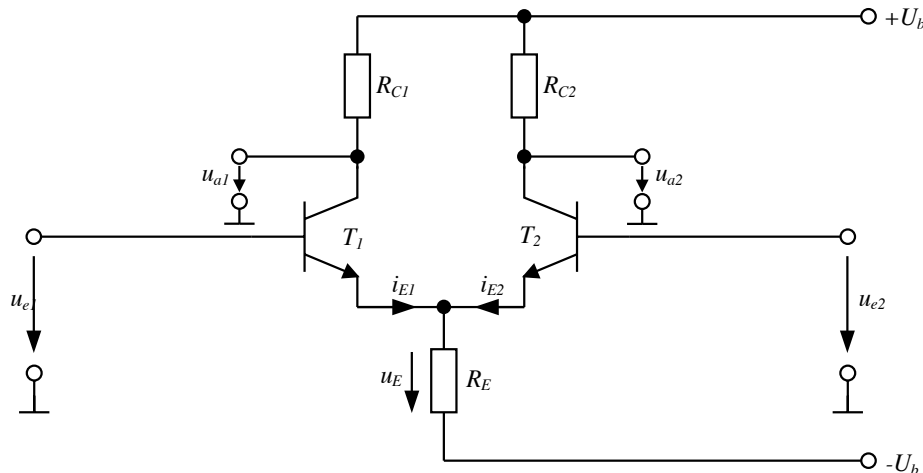


Bild 4.6: Grundsaltung eines Differenzverstärkers.

sehr kleinen Signalen mit äußerst geringer Leistung der Signalquelle sind Störungen, die auf dem Weg vom Signalgeber zum Verstärker in die Verbindungsleitungen einkoppeln können. Für diese Aufgaben sind andere Konzepte für den Schaltungsaufbau erforderlich, als sie bisher in der Vorlesung behandelt wurden.

Eine Lösung bietet hierbei eine Schaltung, wie sie in Bild 4.6 gezeigt wird. Einen Verstärker mit diesem Aufbau nennt man auch Differenzverstärker. Im Folgenden soll durch die Erläuterung der Schaltung die Bezeichnung erklärt und erarbeitet werden. Die Schaltung in Bild 4.6 besitzt zwei Eingänge, die nicht durch Kondensatoren gleichspannungsmäßig von der Signalquelle entkoppelt sind, d.h. die Gleichspannungen an der Basis der beiden Transistoren T_1 und T_2 liegt auch am Innenwiderstand der Quellen an. Damit die gezeigte Schaltung auch einwandfrei funktioniert müssen folgende Bedingungen erfüllt sein:

- beide Transistoren und die Widerstände R_{C1} und R_{C2} müssen jeweils die gleichen Kennwerte besitzen.
- die Transistoren und die Widerstände müssen so angeordnet sein, dass die genannten Paare möglichst auch die gleiche Temperatur besitzen, also möglichst gut thermisch gekoppelt sind.
- für beide Transistoren müssen gleiche Arbeitspunkte eingestellt werden.

Im Gegensatz zu den bisher behandelten Schaltungen besitzt der Verstärker in Bild 4.6 jetzt zwei Eingänge. Wenn wir die Gleichungen der beiden Eingangsflächen erstellen, erhält man:

$$\begin{array}{ll}
 \text{Masche 1} & \text{Masche 2} \\
 u_{e1} - u_E - u_{BE1} = 0 & u_{e2} - u_E - u_{BE2} = 0 \\
 u_{BE1} = u_{e1} - u_E & u_{BE2} = u_{e2} - u_E
 \end{array} \quad (4.30)$$

Die Eingangsspannungen u_{e1} und u_{e2} bestehen aus einem Gleichanteil (Arbeitspunkteinstellung) und einem Wechselanteil (zu verstärkende Wechselspannung). Durch Einsetzen der Steilheit S und Umformung nach $U_{BE} = I_C/S$ werden die Gleichungen der Maschen zu:

$$\begin{aligned} i_{C1} &= (u_{e1} - u_E) \cdot S_1 & i_{C2} &= (u_{e2} - u_E) \cdot S_2 \\ i_E &= i_{E1} + i_{E2} & \text{für } \beta \gg 1 \text{ gilt } i_E &= i_{C1} + i_{C2} \end{aligned} \quad (4.31)$$

Die Spannung u_E am Emitterwiderstand R_E wird durch die Summe der Emitterströme der beiden Transistoren bestimmt und berechnet sich mit Hilfe der Vereinfachung für $\beta \gg 1$ zu:

$$\begin{aligned} u_E &\approx R_E \cdot (i_{C1} + i_{C2}) = R_E \cdot ((u_{e1} - u_E) \cdot S_1 + (u_{e2} - u_E) \cdot S_2) \\ &\approx R_E \cdot S_1 \cdot u_{e1} - R_E \cdot S_1 \cdot u_E + R_E \cdot S_2 \cdot u_{e2} - R_E \cdot S_2 \cdot u_E \\ u_E(1 + R_E \cdot (S_1 + S_2)) &\approx R_E \cdot (u_{e1} \cdot S_1 + u_{e2} \cdot S_2) \\ u_E &\approx \frac{R_E \cdot (u_{e1} \cdot S_1 + u_{e2} \cdot S_2)}{(1 + R_E \cdot (S_1 + S_2))} \end{aligned} \quad (4.32)$$

Da wir vorausgesetzt haben, dass beide Transistoren identische Kennwerte besitzen, müssen also auch die Steilheiten gleich sein. Mit $S_1 = S_2 = S$ wird die Spannung am Emitterwiderstand u_E zu:

$$\begin{aligned} u_E &\approx \frac{S \cdot R_E \cdot (u_{e1} + u_{e2})}{(1 + 2 \cdot S \cdot R_E)} \\ \text{mit } S \cdot R_E &\gg 1 \text{ wird} \\ u_E &\approx \frac{(u_{e1} + u_{e2})}{2} \end{aligned} \quad (4.33)$$

Die Spannung u_E ist proportional der Summe der beiden Eingangsspannungen ($u_{e1} + u_{e2}$), d.h. wenn beide Wechselspannungsanteile die gleiche Polarität und identische Amplituden besitzen, verändert sich u_E proportional der Summe ($u_{e1} + u_{e2}$). Die Basis-Emitterspannung ist nach der Maschengleichung die Differenz der Eingangsspannung und der Spannung am Emitterwiderstand. Verändert sich u_E , werden U_{BE1} und U_{BE2} verändert, wodurch die Arbeitspunkte der beiden Transistoren verschoben werden. Auf das Ausgangskennlinienfeld eines Transistors bezogen, bedeutet dies eine Abnahme bzw. Zunahme der Steuerspannung und damit Abnahme bzw. Erhöhung des Kollektorstromes, was wiederum zu einem reduzierten bzw. erhöhtem Emitterstrom und damit zu einer Verkleinerung bzw. Vergrößerung der Spannung am Emitterwiderstand führt. Diese Art der Gegenkopplung haben wir bisher nur zur Kompensation der Temperatureinflüsse betrachtet. Sie wirkt in dieser Schaltung aber auch, wenn an beiden Eingängen identische Wechselspannungen anliegen. Man spricht in einem solchen Fall auch von Gleichtaktsignalen. Durch die Gegenkopplung ist die Verstärkung dieser Signale klein. Haben die Eingangswechselspannungen dagegen entgegen gesetzte Polarität und gleiche Amplitude

bleibt die Summe der Eingangsspannungen ($u_{e1} + u_{e2}$) konstant, d.h. die Spannung u_E am Emittterwiderstand ändert sich nicht. Damit bleiben auch die Arbeitspunkteinstellungen der beiden Transistoren und die Spannungsverstärkung der Schaltung erhalten, da es keine Gegenkopplung wie in vorher beschriebenen Fall gibt. Diese Betriebsart bezeichnet man als „Gegentaktbetrieb“. Man definiert jetzt ein Gleichtaktsignal

$$u_G = \frac{u_{e1} + u_{e2}}{2} \quad (4.34)$$

und ein Gegentaktsignal

$$u_D = u_{e1} - u_{e2} \quad (4.35)$$

Die Gleichtaktverstärkung wird damit zu:

$$A_G = -\frac{u_{a1}}{u_G} = -\frac{u_{a2}}{u_G} \approx -\frac{R_C}{2 \cdot R_E} \quad (4.36)$$

Beim Gegentaktbetrieb wird die Schaltung in Bild 4.6 jedoch als eine Kombination einer Kollektor- und einer Basisschaltung eingesetzt. Nach Gleichung 3.77 ist die Verstärkung der Kollektorschaltung etwa 1 und nach Gl. 3.86 die der Basisschaltung etwa $S \cdot R_C$. Dies bedeutet, dass das verstärkte Signal von Eingang 1 am Kollektor des Transistors 2 abgegriffen wird und umgekehrt. Die Gegentaktsverstärkung der Schaltung ist

$$\begin{aligned} A_D &= \frac{\Delta u_{a2}}{\Delta u_D} = -\frac{\Delta u_{a1}}{\Delta u_D} = \beta \cdot \frac{R_C \parallel r_{CE}}{2 \cdot r_{BE}} \\ \text{mit } r_{BE} &= \frac{\beta}{S} \text{ und } r_{CE} \gg R_C \text{ wird} \\ A_D &\approx \frac{1}{2} \cdot S \cdot R_C \end{aligned} \quad (4.37)$$

Das Verhältnis der Gegentakts- zur Gleichtaktverstärkung ist

$$G = \frac{|A_D|}{|A_G|} = \frac{|\frac{1}{2} \cdot S \cdot R_C|}{|-\frac{R_C}{2 \cdot R_E}|} = S \cdot R_E \quad (4.38)$$

und wird Gleichtaktunterdrückungsfaktor (Common Mode Rejection Ratio, CMRR) genannt. Verstärker nach Bild 4.6 haben also unterschiedliches Betriebsverhalten, abhängig von der Phasenlage der Eingangsspannungen zueinander. Im so genannten „Gegentaktbetrieb“ (differential mode) ist die Spannungsverstärkung wesentlich höher als im „Gleichtaktbetrieb“ (common mode). Deshalb bezeichnet man diese Art von Verstärkern als Differenzverstärker. Um die Gleichtaktunterdrückung weiter zu verbessern, müsste man den Emittterwiderstand weiter erhöhen. Dies würde aber zu einer Verschiebung des Arbeitspunktes zu immer höheren Gleichanteilen der Eingangsspannungen führen. Deshalb verwendet man anstelle des Emittterwiderstands eine Konstantstromquelle bzw.

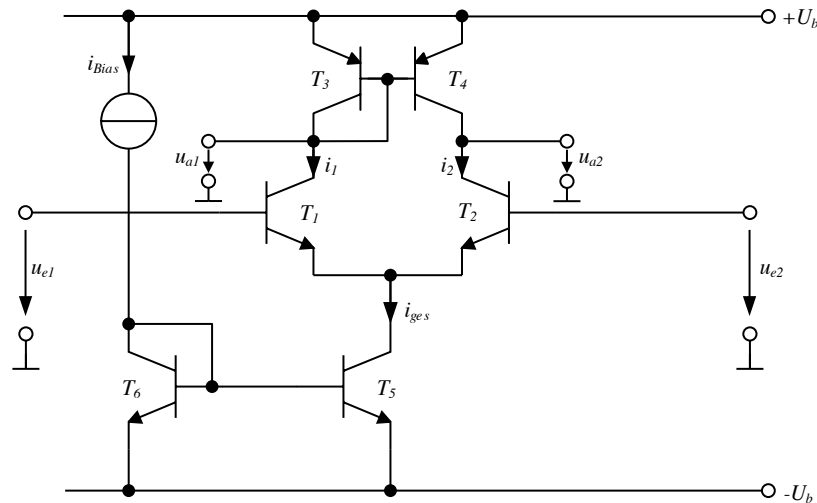


Bild 4.7: Typische Schaltung eines Differenzverstärkers in Bipolartechnologie.

eine Stromspiegelschaltung wie sie im vorangegangenen Kapitel beschrieben wurde. Da die Schaltung auch kleine Gleichspannungen ebenso wie kleine Wechsellspannungen verstärken soll, senkt man das Potential des Arbeitspunktes an der Basis durch Verwendung einer negativen Spannungsquelle soweit ab, dass es gerade dem Massepotential (0 Volt) entspricht. Durch geeignete Dimensionierung der Stromquelle und Ersetzen der Kollektorstufen durch eine Stromspiegelschaltung wird erreicht, dass bei einer Eingangsspannung von 0 V die Ausgangsspannung ebenfalls 0 V ist. Die vollständige Schaltung ist in Bild 4.7 gezeigt. Mit dieser Schaltung werden Gleichtaktunterdrückungsfaktoren bis $G = 10^6$ erreicht. Der Eingangswiderstand eines Differenzverstärkers ist der einer Kollektorschaltung ($\beta \cdot R_E$). Durch den hohen differentiellen Widerstand der Stromquelle kann r_e sehr leicht Werte $> 10^6 \Omega$ annehmen.

4.4 Gegentaktverstärker

Zur Verstärkung von Wechsellspannungen, die sowohl positive als auch negative Amplituden besitzen müssen, insbesondere bei Endstufen von Leistungsverstärkern, werden Gegentaktverstärker eingesetzt. Diese verarbeiten das Eingangssignal in zwei Kanälen. Der eine Kanal verstärkt die positiven und der andere die negativen Anteile des Signals. In der modernen Schaltungstechnik werden dazu vorwiegend komplementäre Kollektorstufen bzw. Drainstufen eingesetzt. Eine einfache Form eines Gegentaktverstärkers ist in Bild 4.8a dargestellt. Der npn Transistor überträgt die positive Halbwelle der sinusförmigen Spannung und ist während der negativen Halbwelle gesperrt. Diese wird vom pnp Transistor übertragen, der während der positiven Halbwelle gesperrt ist. Am Lastwiderstand R_L liegen dann nacheinander beide Halbwellen an. Ist $u_e = 0$, sind beide Transistoren gesperrt, d.h. der Ausgangsstrom ist $i_E = 0$. Eine reale Betrachtung der Ausgangsspannung wird notwendig, da wir bei der Betrachtung der Ausgangskennlinienfelder der Transistoren festgestellt haben, dass unterhalb $U_{BE} = 0,6 - 0,7 \text{ V}$ der

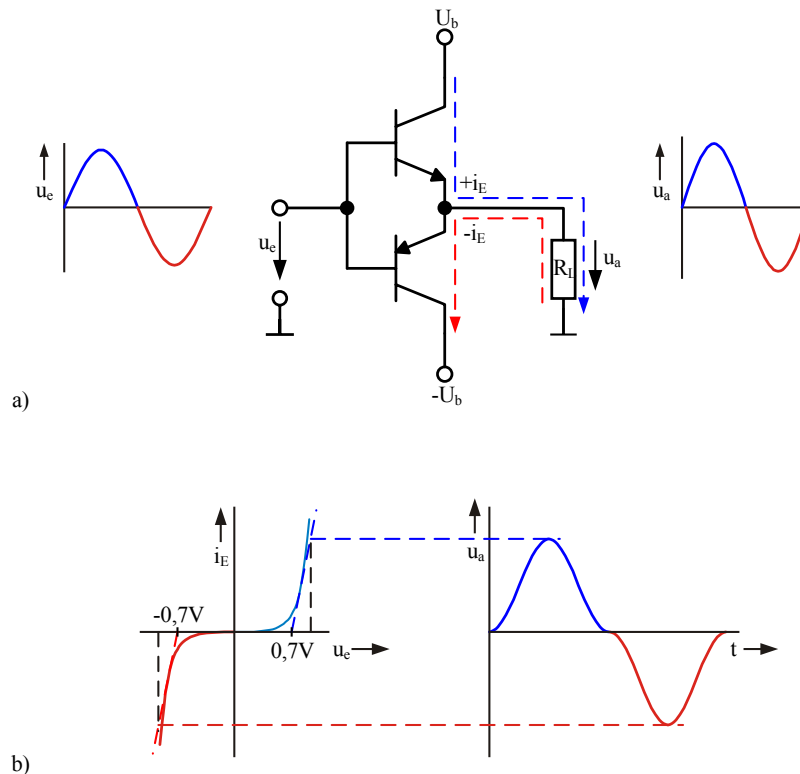


Bild 4.8: a) Kollektorschaltung mit komplementären Transistoren, b) Ansteuerkennlinie und Ausgangsspannung der Gegentaktverstärkerstufe.

Transistor zu sperren beginnt. In Bild 4.8b ist der Zusammenhang zwischen der Eingangsspannung u_e und dem Emittterstrom i_E dargestellt. Man erkennt, dass der Emittterstrom ($i_E = (\beta + 1) \cdot i_B$) unterhalb einer Eingangsspannung von $u_e = 0,7\text{ V}$ sehr stark nichtlinear wird, was dann zu einer starken Verzerrung der Ausgangsspannung im Bereich um 0 V herum führt. Diese Verzerrungen mögen zwar bei sehr großen Ausgangssignalen kaum ins Gewicht fallen, aber bei niedrigen Ausgangsspannungen sind sie doch sehr störend. Diese Verzerrungen können reduziert werden, wenn die Ansteuerung der Schaltung verändert wird.

Bild 4.9a zeigt eine erweiterte Gegentaktstufe. Die Basis der beiden Transistoren wird jetzt über einen Spannungsteiler mit zwei Widerständen und zwei Dioden angesteuert. Wenn jetzt die Eingangsspannung $u_e = 0\text{ V}$ beträgt, liegen an den beiden Basisanschlüssen bereits etwa $0,7\text{ V}$ an, d.h. es fließt bereits ein Ruhestrom durch die Transistoren. Der Verstärker ist damit also auch bei 0 V am Eingang betriebsbereit. Die Ansteuerkennlinie wird dadurch linear, wie in Bild 4.9b dargestellt. Die Basis-Emitttervorspannung für die beiden Transistoren wird durch die Dioden erzeugt. Diese werden über die Widerstände R_1 und R_2 mit dem notwendigen Strom versorgt, so dass jeweils etwa $0,7\text{ V}$ Spannung über den Dioden D_1 und D_2 gemessen werden kann. Der Vorteil des Einsatzes von Dioden liegt in der gleichzeitigen Kompensation der temperaturbedingten

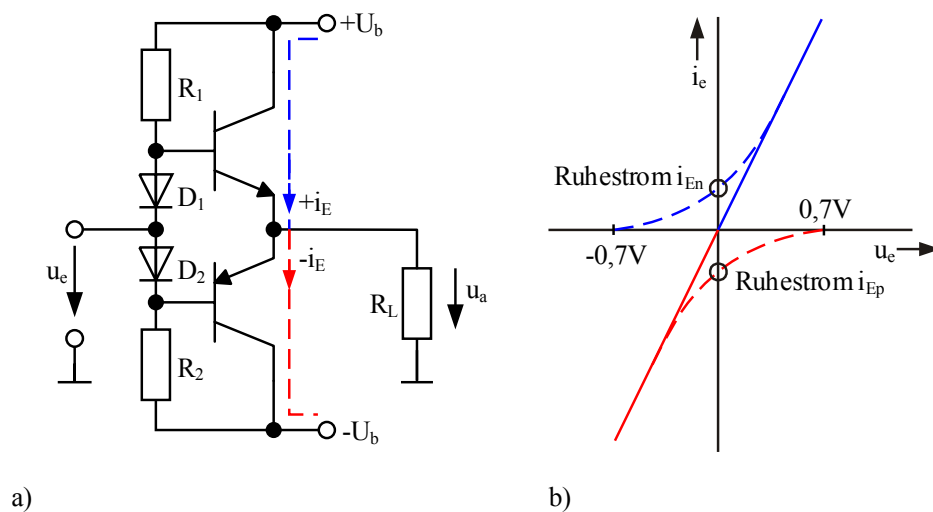


Bild 4.9: a) Erweiterte Kollektorschaltung b) Linearisierte Ansteuerkennlinie der Gegentaktverstärkerstufe.

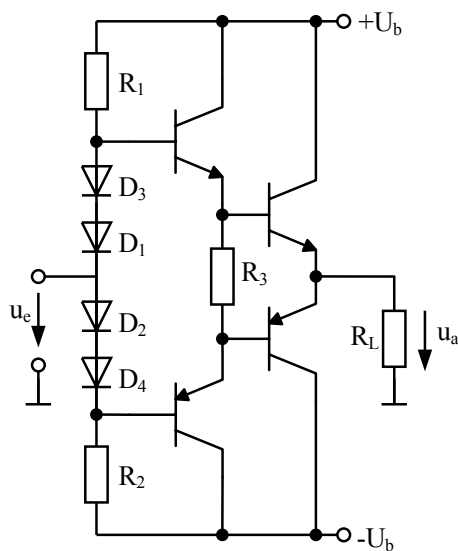


Bild 4.10: Gegentaktverstärkerstufe für hohe Verstärkung.

Schwankungen der Basis–Emitterspannung der Transistoren. Unter den bisher gemachten Voraussetzungen, dass beide Transistoren identische Kennlinienfelder besitzen und an beiden Dioden die gleichen Spannungen abfallen, sind auch die beiden Ruhestrome i_{En} und i_{Ep} gleich und durch die Last fließt kein Strom, da das Ausgangspotential damit auch zu null wird.

Werden Ausgangsstufen mit höherer Verstärkung oder höherer Ausgangsleistung benötigt, muss die Schaltung nach Bild 4.9a um eine weitere Transistorstufe nach Bild 4.10 erweitert werden. Die Transistoren T_3 und T_4 werden dabei zur direkten Ansteuerung der Transistoren T_1 bzw. T_2 eingesetzt. Um eine korrekte Basis–Emittervorspannung zu erzeugen, müssen jetzt jeweils zwei Dioden in Reihe geschaltet werden. Damit wird eine Vorspannung für die beiden jeweils in Reihe geschalteten Basis–Emitterdioden der npn bzw. pnp Transistoren von etwa 1,4 V erzeugt, so dass die Ruhestrome wieder symmetrisch eingestellt werden können. Die Kombination zweier Transistoren wie sie in der Schaltung nach Bild 4.10 eingesetzt wird, nennt man Darlington–Schaltung. Die Stromverstärkung ist dabei das Produkt aus den Einzelstromverstärkungen $\beta_3 \cdot \beta_1$ der npn Transistoren T_1 und T_3 bzw. $\beta_4 \cdot \beta_2$ der pnp Transistoren T_2 und T_4 .

4.5 Rückgekoppelte Verstärker

Bei der Stabilisierung des Arbeitspunktes wurde bei den Transistorgrundschaltungen bereits mehrfach der Begriff „Rückkopplung“ verwendet. Allgemein versteht man darunter, dass ein Teil des Ausgangssignals auf irgendeine Art und Weise auf den Eingang zurückgeführt wird. Das Ausgangssignal kann dabei ein Strom oder eine Spannung sein, bestehend entweder aus einer reinen Gleich– oder Wechselgröße oder einer Kombination aus einer Gleich– und einer Wechselgröße. Bei der Rückkopplung unterscheidet man je nach Verschiebung der Phase zwischen Eingangssignal und dem zurück gekoppelten Signal zwischen „Mitkopplung“ (Phasenverschiebung $n \cdot 2\pi$, mit $n = 0, 1, 2, 3, \dots$) und „Gegenkopplung“ (Phasenverschiebung $(2 \cdot n + 1) \cdot \pi$, mit $n = 0, 1, 2, 3, \dots$). Die Mitkopplung erhöht die Verstärkung. Dies kann besonders bei mehrstufigen Verstärkern dazu führen, dass die Schaltung instabil wird, d.h. die Ausgangssignale werden immer größer, bis sie schließlich an den Aussteuerungsgrenzen enden. Eine lineare Abhängigkeit zwischen Eingangs– und Ausgangsgröße besteht dann nicht mehr. Man spricht auch davon, dass die Schaltung „schwingt“. Diese Instabilität ist bei Verstärkerschaltungen unerwünscht, hingegen bei Oszillatoren, Komparatoren und Schmitt–Triggern für eine einwandfreie Funktion notwendig.

Die Gegenkopplung erniedrigt die Verstärkung. Dies haben wir besonders beim Differenzverstärker und bei der Stabilisierung des Arbeitspunktes durch eine Kollektorstrom–Rückkopplung gesehen. Durch eine Wechselstrom– oder Spannungsgegenkopplung wird aber auch die Verstärkung des Signals reduziert. Dies dient u.a. zur Stabilisierung der Verstärkung von mehrstufigen Verstärkern da damit die Verzerrungen im Großsignalbetrieb reduziert und die Bandbreite der Verstärker erhöht werden kann. Verzerrungen

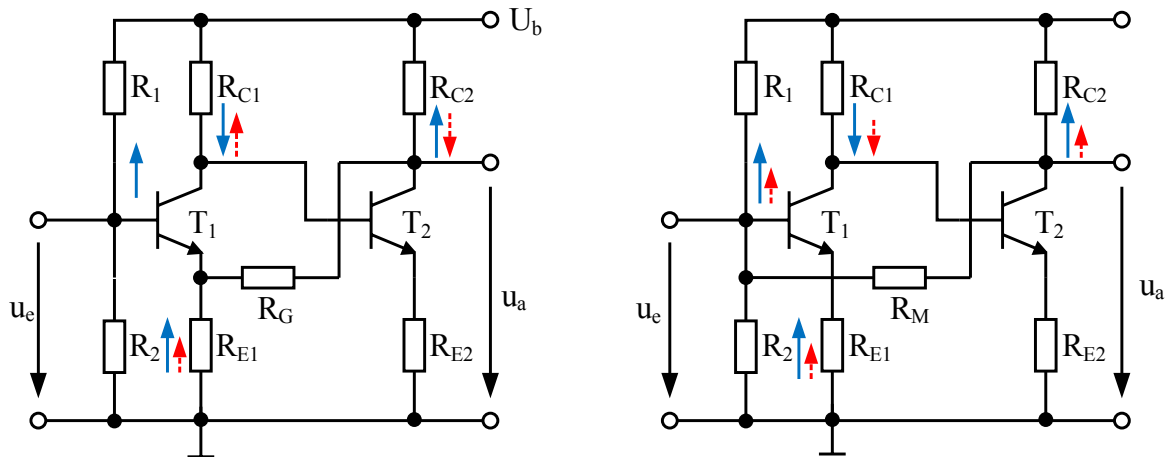


Bild 4.11: Prinzip a) der Gegenkopplung und b) der Mitkopplung für Gleichspannungen.

im Großsignalbetrieb entstehen durch die gekrümmten Kennlinien der Transistoren. Im Kleinsignalbetrieb stellen wir den Arbeitspunkt ein und ändern das Eingangssignal nur wenig, so dass der Kollektorstrom sich auch nur in einem kleinen Bereich verändert. Wir gehen davon aus, dass die Steilheit $S = I_C/U_T$ konstant ist. Im Großsignalbetrieb jedoch ist die Änderung des Kollektorstroms so groß, dass wir diese vereinfachte Annahme nicht aufrechterhalten können. Die Steilheit muss deshalb als veränderlich in Abhängigkeit vom Kollektorstrom betrachtet werden. Bereits bei einstufigen Verstärkern ohne Gegenkopplung gilt:

$$A = -S \cdot R_C$$

während bei einer Verstärkerstufe mit einem zusätzlichen Emittterwiderstand und der Annahme $S \cdot R_E \gg 1$ die Verstärkung unabhängig von S wird:

$$A = -\frac{R_C}{R_E}$$

Die Verstärkung wird also nur noch von der äußeren Beschaltung des Transistors bestimmt. In der Gleichung kommen keine nichtlinearen Größen mehr vor. Am Beispiel der Schaltungen in den Bildern 4.11 und 4.12 sollen die Prinzipien der Gegenkopplung und der Mitkopplung für gleich- bzw. wechsellspannungsgekoppelte zweistufige Verstärker veranschaulicht werden.

4.5.1 Gegenkopplung

Wird die Spannung nach Bild 4.11 an der Basis von Transistor T_1 erhöht, sinkt die Spannung am Kollektor von T_1 und damit auch an der Basis von T_2 ab. Die Spannung

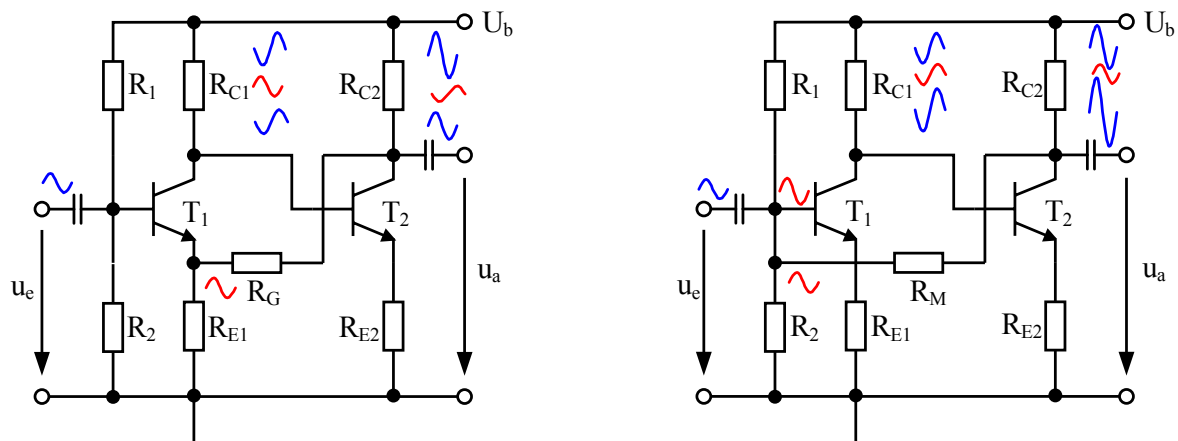


Bild 4.12: Prinzip a) der Gegenkopplung und b) der Mitkopplung für Wechselspannungen.

am Kollektor von T_2 steigt an. Ein kleiner Teil dieser Spannung wird über R_G auf den Emitter von T_1 zurückgeführt. Der Anstieg der Emitterspannung reduziert die Basis–Emitterspannung von T_1 , d.h. die Spannung am Kollektor von T_1 und an der Basis von T_2 wird erhöht. Damit sinkt die Spannung am Kollektor von T_2 . Durch diese Gegenkopplung stabilisiert sich die Ausgangsspannung.

4.5.2 Mitkopplung

Wird die Spannung an der Basis von Transistor T_1 erhöht, sinkt die Spannung am Kollektor von T_1 und damit auch an der Basis von T_2 ab. Die Spannung am Kollektor von T_2 steigt an. Ein kleiner Teil dieser Spannung wird über R_M auf die Basis von T_1 zurückgeführt. Die Erhöhung der Basis–Emitterspannung bewirkt ein Absinken der Spannung am Kollektor von T_1 und an der Basis von T_2 . Damit steigt die Spannung am Kollektor von T_2 weiter an und ebenso die an der Basis von T_1 . Dieser Vorgang wiederholt sich solange, bis die maximale Spannung am Kollektor von T_2 erreicht ist. Durch den Mechanismus der Mitkopplung kann die Schaltung, abhängig vom eingestellten Verstärkungsfaktor schnell die Aussteuergrenze des Verstärkers erreichen. In Bild 4.12a und 4.12b ist die Gegen– und Mitkopplung für Wechselspannungen dargestellt. Der Mechanismus ist derselbe wie bei den Gleichspannungen. Allerdings muss bei einer Mitkopplung an die entstehenden Verzerrungen bei den Ausgangssignalen gedacht werden, da man bei mehrstufigen Verstärkern sehr schnell aus einem Kleinsignalbetrieb in den Großsignalbetrieb kommt. Dann kann auch sehr leicht aus einer sinusförmigen Schwingung am Eingang des Verstärkers eine rechteckförmige Schwingung am Ausgang des Verstärkers entstehen. Die Prinzipien der Gegenkopplung und der Mitkopplung werden insbesondere auch bei Schaltungen mit integrierten Verstärkern, die in einem späteren Kapitel behandelt werden, eine wesentliche Rolle bei der Funktion der aufzubauenden Schaltung spielen.

5 Operationsverstärker

Lernziele

- Aufbau und Wirkungsweise des Operationsverstärkers
- Verstehen der grundlegenden Eigenschaften des Operationsverstärkers
- Kennenlernen und Verstehen der analogen Grundsaltungen mit Operationsverstärkern
- Kennenlernen des Aufbaus von Messverstärkern

5.1 Aufbau und Wirkungsweise

Ein Operationsverstärker ist ein mehrstufiger Verstärker, der als integrierte Schaltung hergestellt wird. Der Verstärker ist so beschaffen, dass seine Wirkungsweise überwiegend durch seine äußere Beschaltung bestimmt werden kann. Um dies zu ermöglichen, werden Operationsverstärker als gleichspannungsgekoppelte Verstärker mit hoher Verstärkung ausgeführt. Damit keine zusätzlichen Maßnahmen zur Arbeitspunkteinstellung erforderlich werden, sind in der Regel zwei Betriebsspannungsquellen erforderlich - eine positive und eine negative. Es gibt heute ein nahezu unüberschaubares Angebot an Operationsverstärkern. Sie unterscheiden sich nicht nur durch ihre Betriebsdaten, sondern auch durch ihren prinzipiellen Aufbau. Anhand des prinzipiellen Aufbaus (Bild 5.1) eines der ersten Operationsverstärker, dem μA 741, soll die Wirkungsweise näher erklärt werden. Da Operationsverstärker universell eingesetzt werden sollen, benötigen sie eine möglichst hohe Differenzverstärkung A_D und einen hohen Gleichtaktunterdrückungsfaktor G . Die geforderten Bedingungen können nur durch den Einsatz mehrerer gekoppelter Verstärkerstufen in einer integrierten Schaltung erreicht werden. Die Transistoren T_1 und T_2 bilden den Eingangsdifferenzverstärker, der in fast jedem Operationsverstärker in irgendeiner Form vorhanden ist. Im Gegensatz zur in Kapitel 4 gezeigten Grundschaltung wird hier nur die Ausgangsspannung am Kollektor von T_2 zur nächsten Stufe weitergeleitet. Der Transistor T_4 ist eine Stromquelle und dient als Arbeitswiderstand für T_2 . Der Strom ist nicht konstant, da T_4 zusammen mit T_3 einen Stromspiegel für den Kollektorstrom I_{C1} darstellt. Damit ergibt sich der Ausgangsstrom der Eingangsstufe zu

$$I_{a1} = I_{C1} - I_{C2} \quad (5.1)$$

Die zusammen geschalteten Emitter des Differenzverstärkers werden über eine Konstantstromquelle, die hier ebenso wie die Stromquelle I_2 schaltungstechnisch nicht aufgelöst

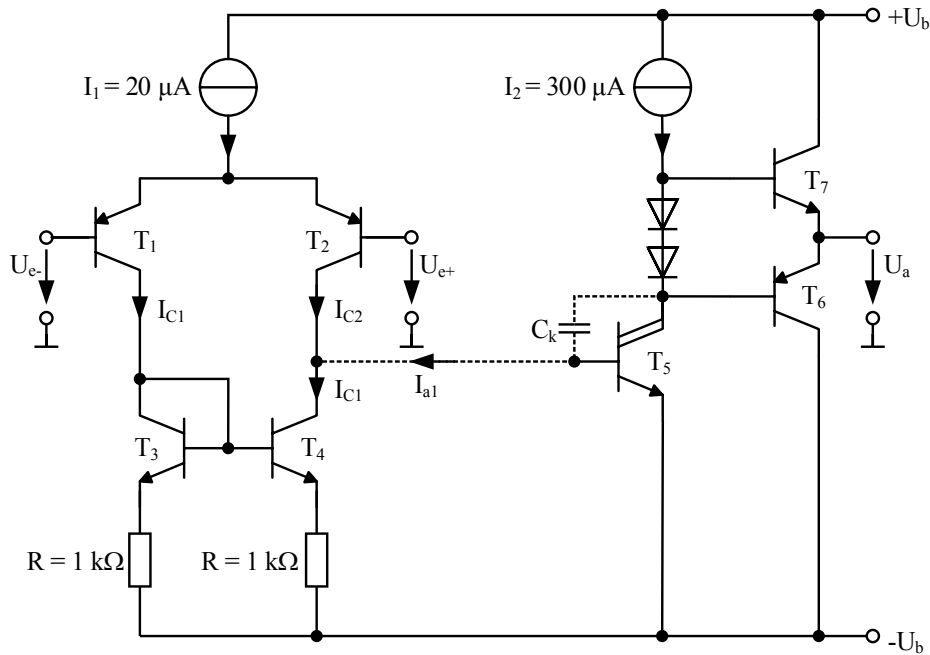


Bild 5.1: Prinzipieller Aufbau des Operationsverstärkers μA 741.

ist, mit der positiven Betriebsspannung verbunden. Führt man die Emitteranschlüsse von T_3 und T_4 nach außen, anstelle sie wie in Bild 5.1 intern mit der negativen Betriebsspannung zu verbinden, kann man die Kollektorruhestrome der beiden Transistoren durch ein Potentiometer, dessen Mittenabgriff dann mit $-U_b$ verbunden wird, gegensinnig verändern und damit eine Nullpunkteinstellung vornehmen. Der Transistor T_5 bildet eine weitere Verstärkerstufe, der auch noch eine zusätzliche Stufe entlang der gestrichelt gezeichneten Verbindungslinie vorgeschaltet sein kann. Das Symbol zeigt, dass es sich hierbei um eine Darlington-Stufe handelt, die beim Gegentaktverstärker im vorigen Kapitel gezeigt wurde. T_5 wird in einer Emitterschaltung mit einer Konstantstromquelle anstelle eines einfachen Lastwiderstands betrieben. Die Transistoren T_6 und T_7 bilden die Ausgangsstufe. Sie arbeiten als komplementäre Emitterfolger mit sehr kleinem Ruhestrom im Gegentaktbetrieb (Beschreibung in Kapitel 4). Die Gesamtverstärkung der Schaltung ist das Produkt der Einzelverstärkungen der hintereinander geschalteten Verstärkerstufen. Durch den Einsatz eines Differenzverstärkers als Eingangsstufe werden Gleichtaktsignalen unterdrückt und Differenzsignale hoch verstärkt. Es können Differenzverstärkungen im Bereich $A_D = 10^5 \dots 10^6$ erreicht werden. Dadurch ist eine starke Gegenkopplung über eine äußere Beschaltung möglich. Die Ausgangsspannung des Operationsverstärkers

$$U_a = A_D \cdot U_D = A_D \cdot (U_{e+} - U_{e-}) \quad (5.2)$$

ist also gleich der verstärkten Eingangsspannungsdifferenz U_D . Der gestrichelt eingezeichnete Kondensator C_k zwischen Kollektor und Basis von T_5 dient zur Korrektur des

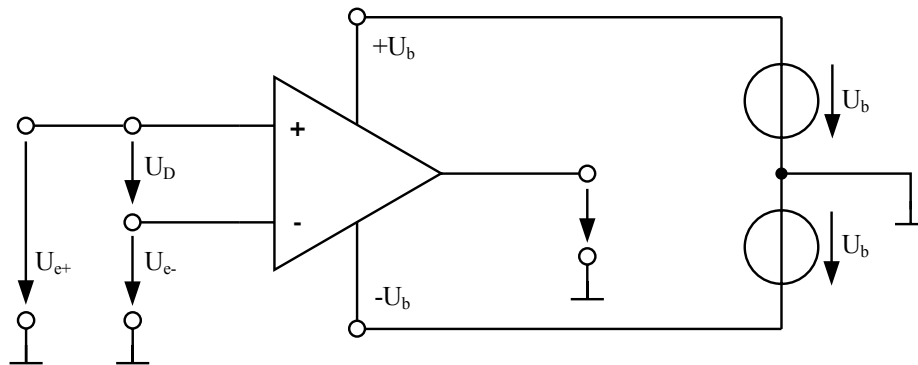


Bild 5.2: Anschlussbelegung eines Operationsverstärkers.

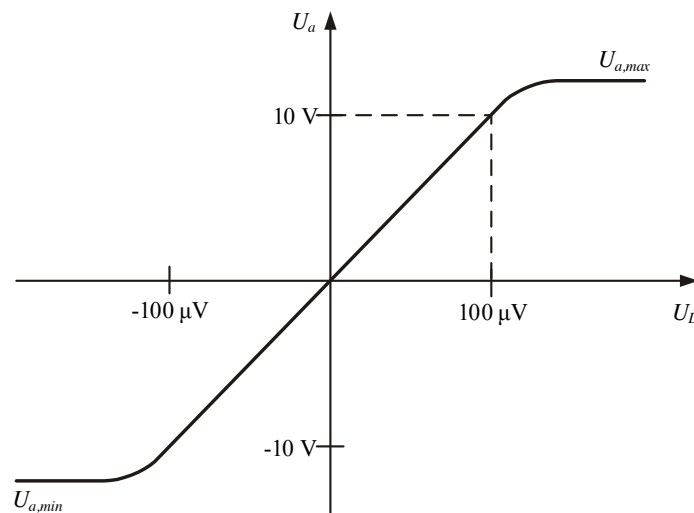


Bild 5.3: Übertragungskennlinie eines Operationsverstärkers.

Frequenzgangs der Schaltung. Auf seinen Einfluss auf das Verhalten der Schaltung soll an dieser Stelle nicht näher eingegangen werden.

Bild 5.2 zeigt das Schaltsymbol von Operationsverstärkern. Die beiden Eingänge werden als invertierender ($-$) und nichtinvertierender ($+$) Eingang bezeichnet. Zur Versorgung besitzt der Operationsverstärker zwei Betriebsspannungsanschlüsse, an die eine positive und negative Betriebsspannung angelegt wird. Operationsverstärker besitzen selbst keinen Masseanschluß, obwohl die Eingangs- und Ausgangsspannungen darauf bezogen werden. Übliche Betriebsspannungen sind $\pm 15\text{ V}$ für Universalanwendungen. Heute werden vermehrt auch Operationsverstärker mit Betriebsspannungen von $\pm 5\text{ V}$ eingesetzt und der Trend geht zu einer weiteren Verkleinerung der Betriebsspannung.

Die Übertragungskennlinie eines realen Operationsverstärkers ist in Bild 5.3 dargestellt. Die Differenzverstärkung ist gegeben durch

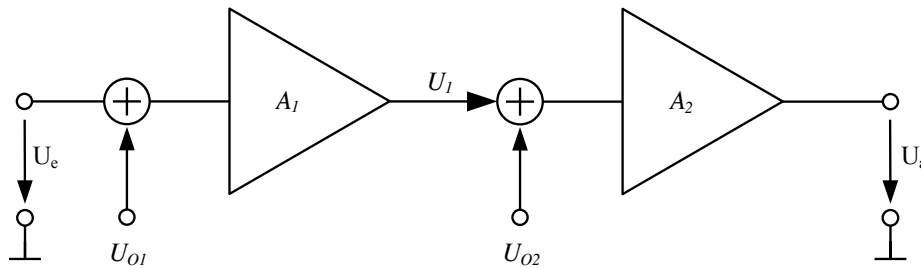


Bild 5.4: Modell zum Einfluss der Offsetspannung in mehrstufigen Verstärkern.

$$A_D = \left. \frac{\partial U_a}{\partial U_D} \right|_{AP} \quad (5.3)$$

Sie entspricht der Steigung im Bild 5.3. Man sieht, dass Bruchteile von einigen mV ausreichen, um den Ausgang voll auszusteuern. Der lineare Arbeitsbereich wird oft auch als Ausgangsaussteuerbarkeit bezeichnet. Wenn die Aussteuerergrenze erreicht ist, steigt U_a bei weiterer Vergrößerung von U_D nicht weiter an, d.h. der Verstärker wird übersteuert. Fassen wir die bisherigen Beschreibungen zusammen, erhält man für einen idealen Operationsverstärker folgende Aussagen:

- der Eingangswiderstand r_e der beiden Eingänge ist sehr hoch ($r_e \rightarrow \infty$)
- der Ausgangswiderstand r_a der Gegentaktendstufe ist sehr klein ($r_a \rightarrow 0$)
- die Differenzverstärkung A_D im Leerlauf ist sehr groß ($A_D \rightarrow \infty$)

Da integrierte Operationsverstärker in großen Stückzahlen hergestellt werden, gibt es bei der Herstellung durch die Vielzahl der Herstellungsschritte auch Toleranzen bei den Bauelementen. Diese führen dazu, dass die Kenngrößen der Operationsverstärker schwanken können und Eingriffe des Anwenders zur Optimierung seiner Schaltung notwendig machen. Die Hersteller geben deshalb nur typische Werte bzw. min. und max. Werte für die einzelnen Kenngrößen an und führen zusätzliche Anschlüsse zur Beschaltung durch den Anwender nach außen. Im Wesentlichen sind das Anschlüsse zur Kompensation von Verschiebungen der Eingangsruhestrome und zur Kompensation des Frequenzgangs des Verstärkers. Im ersten Fall spricht man auch von Offsetspannungs- bzw. Offsetstromkompensation. Zum besseren Verständnis soll zunächst der Begriff „Offsetspannung“ (U_O) erläutert werden. Bild 5.4 zeigt ein Modell hierfür. Allgemein ist die Ausgangsspannung

$$U_a = A_D \cdot (U_D - U_O) \quad (5.4)$$

Setzt man umgekehrt die Bedingung, die Ausgangsspannung soll null sein, und nimmt als Beispiel zwei Verstärkerstufen nach Bild 5.4 mit verschiedenen Offsetspannungen an,

muss dazu als Eingangsspannung die Offsetspannung

$$U_e(U_a = 0) = U_O = -U_{O1} - \frac{1}{A_1} \cdot U_{O2} \quad (5.5)$$

angelegt werden. So gesehen ist die Offsetspannung also die Spannung, die am Eingang eines Operationsverstärkers angelegt werden muss, damit die Ausgangsspannung null wird. Diese allgemein gültigen Betrachtungen über Aufbau, Eigenschaften und Wirkungsweise von Operationsverstärkern gelten zunächst nur für das Bauelement ohne zusätzliche äußere Beschaltung. Der Vorteil der Operationsverstärker liegt jedoch in den vielfältigen Möglichkeiten, durch äußere Beschaltung des Verstärkers die Eigenschaften der Schaltung den gewünschten Anforderungen anzupassen. Deshalb sollen als nächstes zwei Grundschaltungen des Operationsverstärkers untersucht werden.

5.2 Grundschaltungen mit Operationsverstärkern

5.2.1 Der invertierende Verstärker

Die Schaltung des invertierenden Verstärkers ist in Bild 5.5 gezeigt. Die Eingangsspannung U_e wird über den Widerstand R_I an den invertierenden Eingang angeschlossen. Ein Teil der Ausgangsspannung U_a wird über den Widerstand R_N auf den selben Eingang zurückgeführt. Diese Rückkopplung vom Ausgang zum invertierenden Eingang des Verstärkers stellt eine Gegenkopplung dar. Da wir den invertierenden Eingang als virtuellen Massepunkt betrachten können, ergibt sich für den invertierenden Verstärker aus dem Knotensatz folgendes:

$$\frac{U_e}{R_I} + \frac{U_a}{R_N} = 0 \quad (5.6)$$

Hierbei haben wir die wichtigste Regel zur Berechnung von Verstärkerschaltungen mit Operationsverstärkern benutzt. Sie lautet:

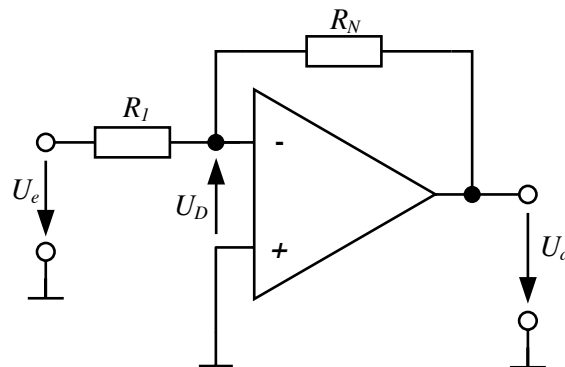


Bild 5.5: Schaltung eines invertierenden Verstärkers.

Die Ausgangsspannung eines beschalteten Operationsverstärkers stellt sich so ein, dass die Eingangsspannungsdifferenz gleich Null wird.

Voraussetzung dafür ist, dass die Schleifenverstärkung groß genug ist, und dass wirklich Gegenkopplung und keine Mitkopplung vorliegt. Somit ergibt sich für die Verstärkung des invertierenden Verstärkers:

$$A = \frac{U_a}{U_e} = -\frac{R_N}{R_1} \quad (5.7)$$

Die Verstärkung der Schaltung wird also nur noch durch die beiden äußeren Widerstände bestimmt und nicht durch die Eigenschaften des Operationsverstärkers.

5.2.2 Der nichtinvertierende Verstärker

Bild 5.6 zeigt die Schaltung des nichtinvertierenden Verstärkers. Betrachten wir den Verstärkungsprozess qualitativ. Lassen wir die Eingangsspannung von null auf einen positiven Wert U_e ansteigen. Im ersten Augenblick ist die Ausgangsspannung noch null und damit auch die rückgekoppelte Spannung. Dadurch tritt am Verstärkereingang die Spannung $U_D = U_e$ auf. Da diese Spannung mit der hohen Differenzverstärkung A_D verstärkt wird, steigt U_a schnell auf positive Werte an und damit auch die rückgekoppelte Spannung am invertierenden Eingang. Dadurch verkleinert sich U_D . Die Tatsache, dass die Ausgangsspannungsänderung der Eingangsspannungsänderungen entgegenwirkt, ist typisch für die Gegenkopplung. Man kann daraus folgern, dass sich ein stabiler Endzustand einstellen wird. Zur quantitativen Berechnung des eingeschwungenen Zustandes geht man davon aus, dass die Ausgangsspannung soweit ansteigt, dass sie gleich der verstärkten Eingangsspannungsdifferenz ist. Die Verstärkung lässt sich wieder leicht durch den Knotensatz berechnen:

$$A = \frac{U_a}{U_e} = 1 + \frac{R_N}{R_I} \quad (5.8)$$

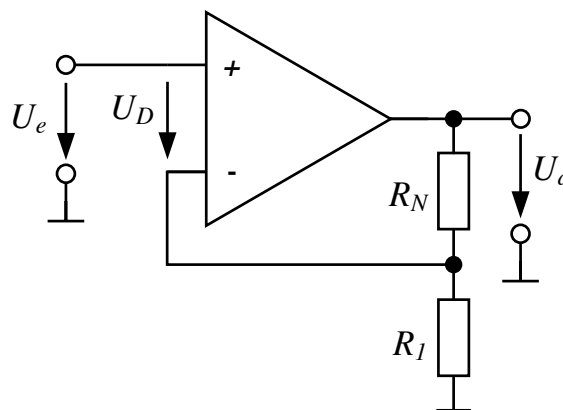


Bild 5.6: Schaltung eines nichtinvertierenden Verstärkers.

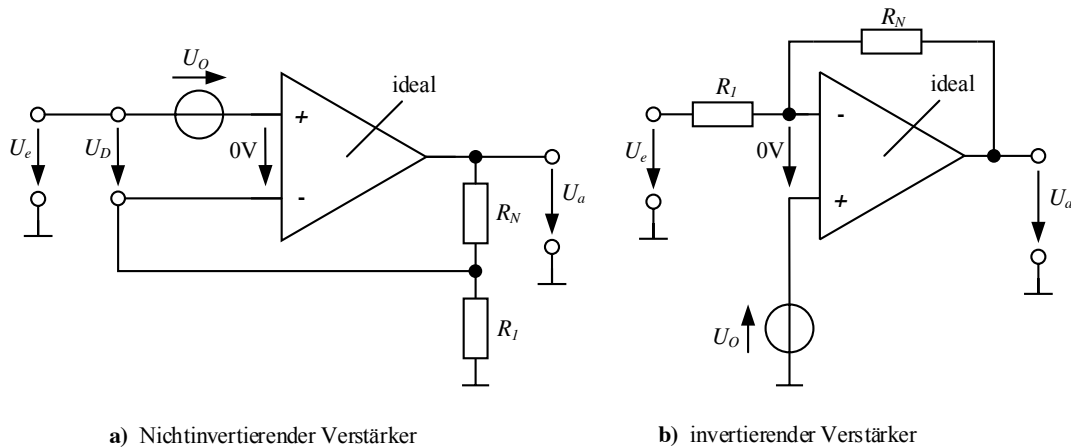


Bild 5.7: Einfluss der Offsetspannung auf den nichtinvertierenden und invertierenden Verstärker.

5.3 Offsetspannung und Offsetstrom

Die im vorigen Kapitel besprochene Offsetspannung wirkt sich auf die beiden gerade behandelten Verstärkerschaltungen aus. Um die Auswirkung der Offsetspannung auf diese stark gegengekoppelten Schaltungen näher zu untersuchen, kann man sich Ersatzschaltbilder nach Bild 5.7 erstellen. Wir betrachten dabei den Operationsverstärker als ideal. Setzt man die Eingangsspannung U_e zu null, sind die Bedingungen für beide Schaltungen gleich. Am Ausgang des nichtinvertierenden Verstärkers erhält man dann wie bereits besprochen die verstärkte Offsetspannung:

$$U_a(U_e = 0) = - \left(1 + \frac{R_N}{R_1} \right) \cdot U_O \quad (5.9)$$

Die Offsetspannung wird also mit dem gleichen Faktor wie das Eingangssignal verstärkt und nicht mehr mit der sehr hohen Leerlaufverstärkung. Wenn wir beim invertierenden Verstärker wie in Bild 5.7b gezeigt, die Offsetspannung am nichtinvertierenden Eingang anlegen, gilt Gleichung 5.9 auch mit ausreichender Näherung für diese Schaltung. Nachdem wir bisher die Operationsverstärker mit Ausnahme der Offsetbetrachtung als „ideal“ angenommen haben, sollen im Folgenden einige reale Eigenschaften der Operationsverstärker beschrieben werden. Dies sind: Eingangsströme, Eingangs- und Ausgangswiderstände. Die Eingangsströme an den beiden Eingängen eines Operationsverstärkers entsprechen dem Basisstrom (bzw. dem Gatestrom bei FETs) der Eingangstransistoren. Seine Größe hängt vom verwendeten Transistortyp und von der Eingangsschaltung ab. Im breiten Angebot der verschiedenen Typen von Operationsverstärkern liegen die Werte der Eingangsströme im Bereich von wenigen pA bis hin zu einigen μA . Da die Eingangsschaltung wie gezeigt mit einem konstanten Kollektorstrom betrieben wird, sind auch die Basisströme praktisch gleich groß. Um Parameterschwankungen der Bauteile (wie auch bei der Offsetspannung) zu beschreiben, werden im Datenblatt der Verstärker

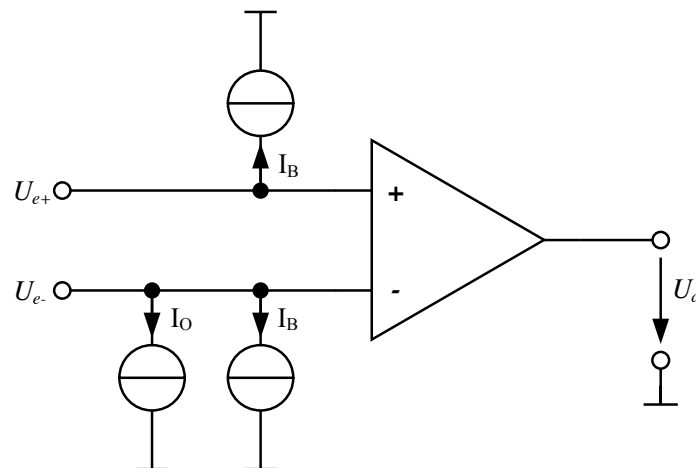


Bild 5.8: Vereinfachte Darstellung der Eingangsruhestrome und des Offsetstroms.

ein mittlerer Eingangsruhestrom (input bias current)

$$I_B = \frac{1}{2} \cdot (I_{e+} - I_{e-}) \quad (5.10)$$

und ein Offsetstrom (input offset current)

$$I_O = |I_{e+} - I_{e-}| \quad (5.11)$$

angegeben. Aus diesen Angaben kann der Anwender leicht die Eingangsströme des Operationsverstärkers berechnen:

$$\begin{aligned} I_{e+} &= I_B \pm \frac{I_O}{2} \\ I_{e-} &= I_B \mp \frac{I_O}{2} \end{aligned} \quad (5.12)$$

Danach werden jedem Eingang ein Ruhestrom und die Hälfte des Eingangs-Offsetstroms zugeordnet. Zur Vereinfachung kann man aber auch den Offsetstrom nur einem Eingang zuordnen, wie dies in Bild 5.8 gezeigt ist. Der dadurch bedingte mögliche Fehler ist sehr klein, denn in der Regel ist der Offsetstrom sehr viel kleiner als der Ruhestrom. Die Wirkung der Eingangsruhestrome und der Eingangsoffsetstroms auf die beiden bisher behandelten Verstärkerschaltungen kann mit Hilfe von Bild 5.9 berechnet werden. Mit der vereinfachten Darstellung der Eingangsströme nach Bild 5.9 und der Bedingung, dass die Widerstände nach folgender Regel

$$R_g = \frac{R_1 \cdot R_N}{(R_1 + R_N)} \quad (5.13)$$

dimensioniert sind, wird die Ausgangsspannung des nichtinvertierenden Verstärkers zu:

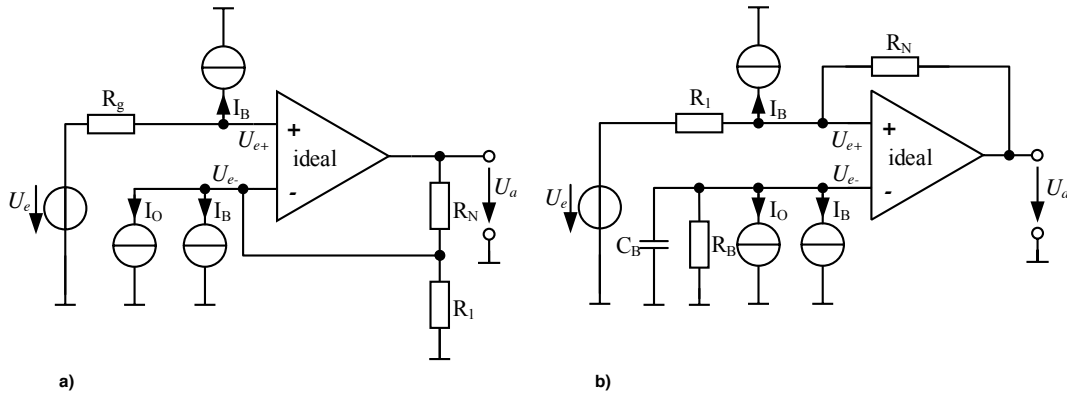


Bild 5.9: Wirkung der Eingangsströme beim a) nichtinvertierenden und b) invertierenden Verstärker.

$$U_a = \left(1 + \frac{R_N}{R_1}\right) \cdot U_e + I_O \cdot R_N \quad (5.14)$$

Der Abgleich von R_g nach der genannten Bedingung bewirkt, dass die an den Eingängen des Operationsverstärkers liegenden Widerstände gleich groß sind. Der Eingangsruhestrom I_B hat damit keinen Einfluss mehr auf die Ausgangsspannung U_a . Es bleibt nur noch ein Fehler durch den Offsetstrom I_O . Da dieser aber meist sehr klein ist, lohnt es sich nicht, einen sehr aufwendigen Abgleich vorzunehmen, wie dies für die Eingangsoffsetspannung üblich ist. Beim invertierenden Verstärker wird an den nichtinvertierenden Eingang, der normalerweise direkt an Masse gelegt wird, ein Widerstand R_B angeschlossen, dessen Wert nach der Bedingung

$$R_B = \frac{R_1 \cdot R_N}{(R_1 + R_N)} \quad (5.15)$$

ermittelt werden kann. Damit sind wie beim nichtinvertierenden Verstärker die Gesamtwiderstände an beiden Eingängen gleich. Für die Ausgangsspannung ergibt sich:

$$U_a = -\frac{R_N}{R_1} \cdot U_e + I_O \cdot R_N \quad (5.16)$$

Damit bleibt auch hier der gleiche durch den Offsetstrom bedingte Fehler wie beim nichtinvertierenden Verstärker ($I_O \cdot R_N$). Zum besseren Verständnis der gerade eingesetzten Regeln für die Dimensionierung der Widerstände R_g und R_B sollen die realen Eingangswiderstände des Operationsverstärkers genauer untersucht werden.

5.4 Eingangs- und Ausgangswiderstand

Zur Veranschaulichung der Wirkung des Eingangs- und des Ausgangswiderstandes

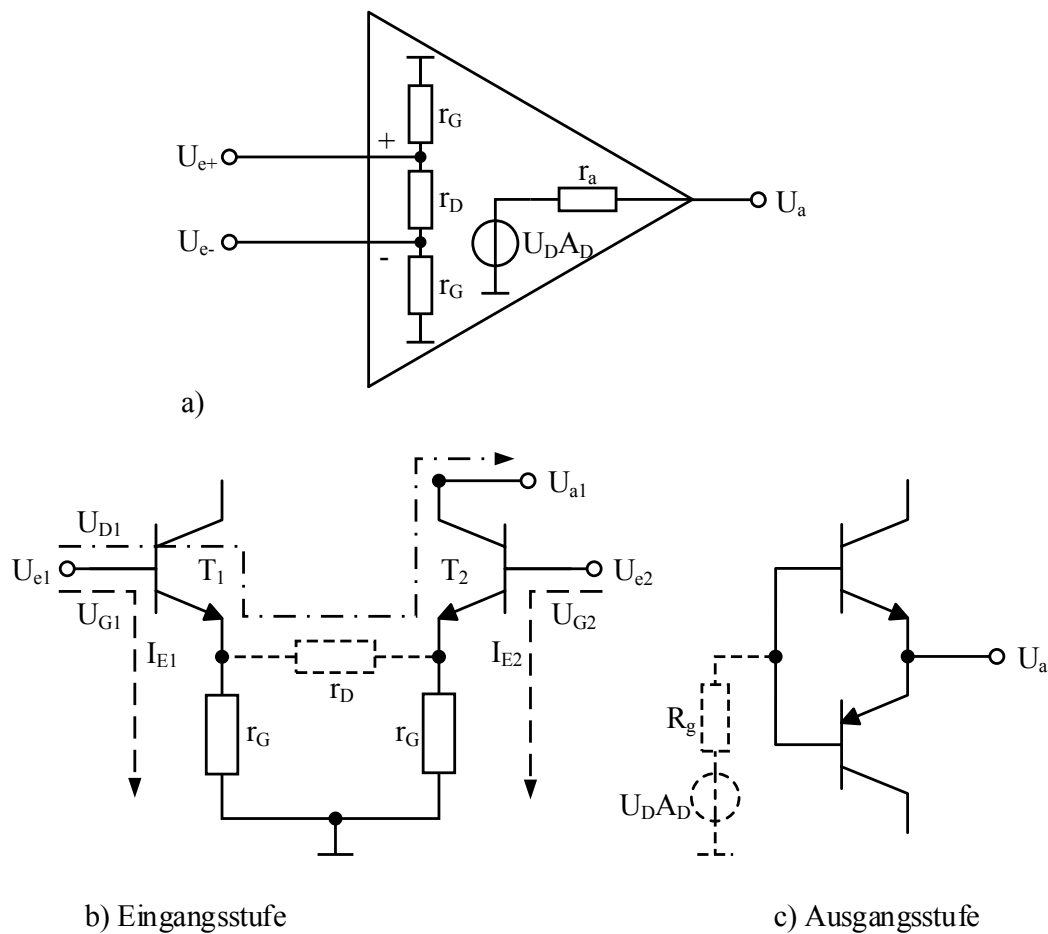


Bild 5.10: Gegentakt- und Gleichtakteingangswiderstand und Ausgangswiderstand beim Operationsverstärker.

soll Bild 5.10 dienen. In Bild 5.10a sind die beim Operationsverstärker wirksamen Widerstände schematisch eingezeichnet. Da der Eingangsverstärker eines Operationsverstärkers ein Differenzverstärker ist, kann man nach Bild 5.10b zwei Eingangswiderstände unterscheiden: den Differenz-Eingangswiderstand und den Gleichtakt-Eingangswiderstand. Der Gleichtakt-Eingangswiderstand wird bestimmt durch den Eingangswiderstand einer Emitterschaltung mit Stromgegenkopplung $r_e = r_{BE} + \beta \cdot R_E$ mit sehr großem Emitterwiderstand R_E , der dem Widerstand der Konstantstromquelle entspricht. Damit können sehr hohe Widerstandswerte erreicht werden. Die Gleichtakteingangswiderstände r_G liegen zwischen dem jeweiligen Eingang und Masse, also parallel zu den Eingängen. Zur Untersuchung des Differenzeingangswiderstandes muss man den Weg des Gegentaktsignals näher betrachten. Die Eingangsstufe wirkt jetzt wie eine Reihenschaltung einer Kollektor- und einer Basisschaltung. Aus der Betrachtung der Eingangswiderstände der Grundschaltungen in Kapitel 3 wissen wir, dass der Eingangswiderstand einer Kollektorschaltung $r_e = \beta \cdot R_E$ ist, und der einer Basisschaltung etwa $1/S$. Damit wird der Differenzeingangswiderstand r_D durch den hohen Eingangswiderstand der Kollektor-

schaltung bestimmt. Der Ausgangswiderstand wird durch die Gegentaktendstufe bestimmt, die wiederum eine Kollektorschaltung darstellt. Wenn kein Emitterwiderstand vorhanden ist, wird $r_a = (R_g/\beta + 1/S)$, wobei R_g der Ausgangswiderstand der vorgeschalteten Verstärkerstufe im Operationsverstärker ist. Die Spannung der Quelle ist die mit der Differenzverstärkung A_D verstärkte Differenz der Eingangsspannungen. Der Ausgangswiderstand r_a kann damit sehr kleine Werte annehmen.

5.5 Frequenzgang

Beim Studium des Datenblatts eines Operationsverstärkers findet man neben einer Fülle von technischen Daten eine Vielzahl von graphischen Darstellungen, die Auskunft über das Verhalten des Bauelements geben. Eine wesentliche Angabe ist dabei die Darstellung der Differenzverstärkung A_D (Open-Loop Signal Differential Voltage Amplification) über der Frequenz, wie sie für den schon erwähnten Operationsverstärker $\mu A 741$ in Bild 5.11 gezeigt ist. Open-Loop bedeutet, dass der Operationsverstärker ohne äußere Gegenkopplung betrachtet wird. Die Darstellung erfolgt üblicherweise in dB (Dezibel) über dem Logarithmus der Frequenz, wobei

$$A_D \Big|_{dB} = 20 \cdot \log \frac{u_a}{u_e} \quad (5.17)$$

ist. Dies bedeutet: 20 dB entsprechen einer Spannungsverstärkung von 10, 40 dB einer von 100 usw.

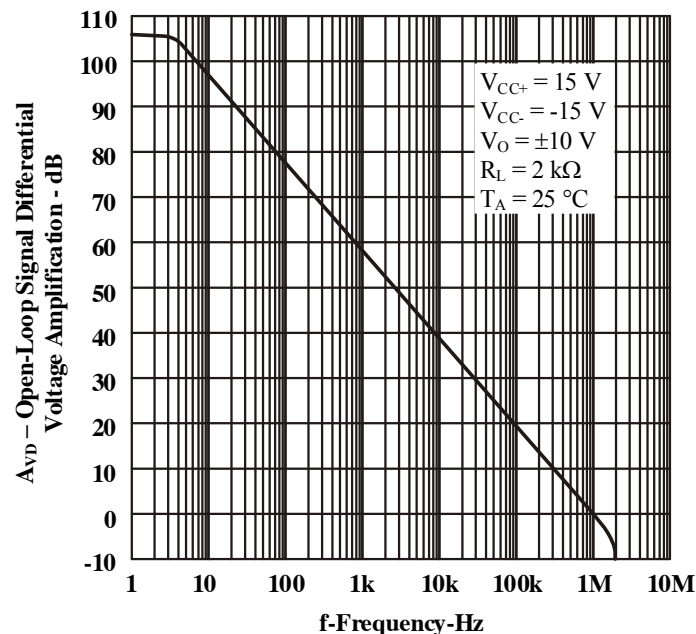


Bild 5.11: Frequenzgang der Differenzverstärkung A_D eines Operationsverstärkers $\mu A 741$.

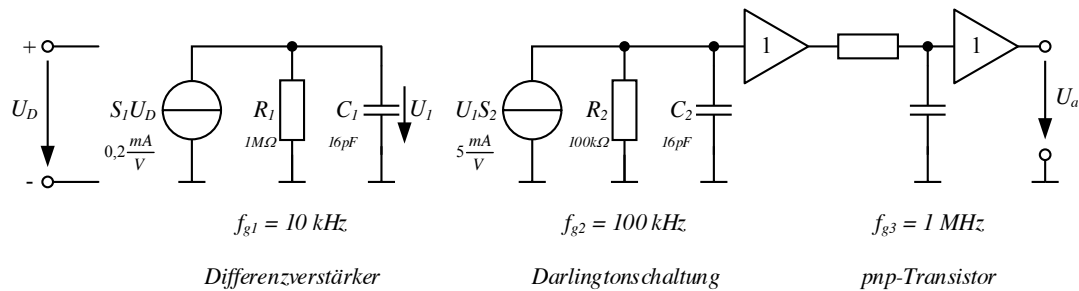


Bild 5.12: Grenzfrequenzen in einem Operationsverstärker nach Bild 5.1.

Auf den Verstärkungsfaktor bezogen, bedeutet dies eine doppellogarithmische Darstellung, aber wie Bild 5.11 zeigt, wird daraus bei der Angabe in dB eine einfach logarithmische.

Man kann aus der Darstellung folgendes erkennen: Bei Frequenzen unterhalb 3 Hz beträgt A_D über 100 dB, ist also größer als 10^5 . Danach wird mit steigender Frequenz A_D immer kleiner, bis bei etwa 1 MHz die Verstärkung auf 0 dB bzw. den Wert 1 abgesunken ist. Die Steigung der Geraden beträgt etwa -20 dB pro Dekade der Frequenz. Dies entspricht einem Tiefpass. Für den Anwender lässt sich daraus die maximal mögliche Frequenz ablesen, die für eine gewünschte Verstärkung einer Eingangsspannung möglich ist. Beachten sollte man auch die Randbedingungen, unter denen der gezeigte Frequenzgang gilt. Diese sind hier im Bild direkt mit angegeben: Betriebsspannung $\pm 15 \text{ V}$, Ausgangsspannung $\pm 10 \text{ V}$, Lastwiderstand am Ausgang $2 \text{ k}\Omega$ und 25° C Umgebungstemperatur. Neben dem sogenannten Amplitudengang ist aber für die einwandfreie Funktion einer Schaltung auch der sogenannte Phasengang entscheidend. Beim Phasengang wird die Phasenverschiebung des Ausgangssignals gegenüber dem Eingangssignal dargestellt. Die Phasenverschiebung ist bei niederfrequenten Signalen praktisch ausschließlich durch parasitäre Kapazitäten der Bauelemente verursacht, während bei hoch- und höchstfrequenten Verstärkern auch die Laufzeit des Signals durch den Verstärker mit zur Phasenverschiebung beiträgt.

Die Stabilität einer Verstärkerschaltung hängt von der Phasenlage zwischen Ausgangs- und Eingangssignal ab. Da Operationsverstärker immer mehrstufige Verstärker sind, und jede einzelne Stufe ein eigenes Tiefpassverhalten zeigt, ist es sinnvoll, sich ein Modell mit den wichtigsten Tiefpässen zu erstellen. Für die Schaltung in Bild 5.1 würde dies der Darstellung in Bild 5.12 entsprechen. Die Eingangsstufe besitzt wegen der geringen Eingangsströme und der hohen Eingangswiderstände die niedrigste Grenzfrequenz. Mit den in Bild 5.12 angenommenen Werten beträgt diese 10 kHz. Die Grenzfrequenz der Darlingtonstufe ist aufgrund der wesentlich höheren Ströme und kleineren wirksamen Widerstände schon deutlich höher. Die Grenzfrequenz der Gegentaktendstufe wird im Wesentlichen durch den pnp Transistor bestimmt, der bei sehr preiswerten Technologien oft schlechtere Daten als der npn Transistor besitzt. Prinzipiell ist mit jedem Tiefpass oberhalb der Grenzfrequenz eine Abnahme der Verstärkung um 20 dB/Dekade und eine

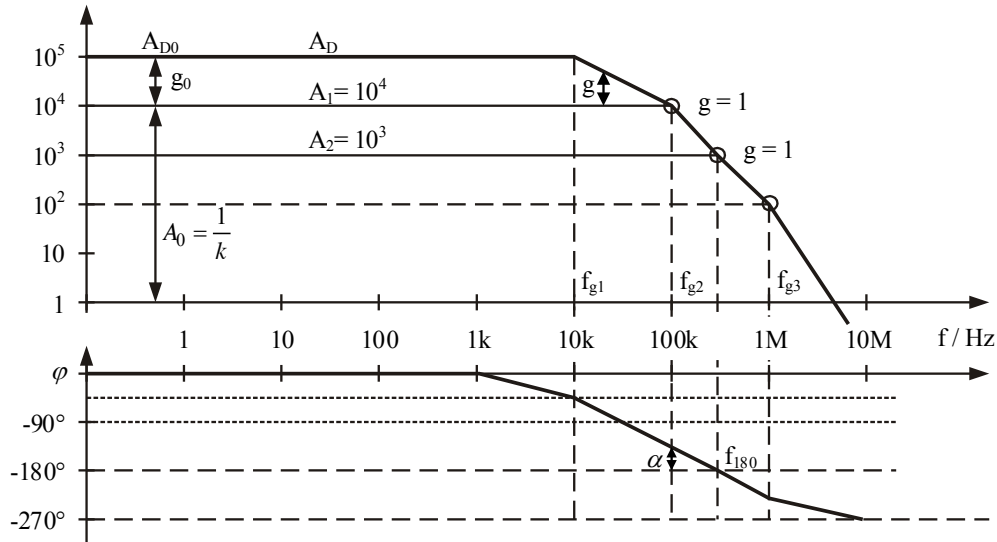


Bild 5.13: Bode-Diagramm für die Ersatzschaltung nach Bild 5.12.

zusätzliche Phasenverschiebung verbunden. Bei der Grenzfrequenz beträgt die Phase 45° und wächst bei höher werdenden Frequenzen bis auf 90° an. Aus den Überlegungen und Berechnungen nach Bild 5.12 kann ein Bode-Diagramm aus Geradenstücken konstruiert werden, wie es in Bild 5.13 dargestellt ist. Man erkennt bei f_{g1} ein Absinken des Betrags der Verstärkung A_D um 20 dB/Dekade und bereits eine Verschiebung der Phase um 45° . Diese verändert sich nun mit $-90^\circ/\text{Dekade}$, so dass sie bei $f_{g2} = 100 \text{ kHz}$ bereits 135° beträgt. Ab hier sinkt A_D mit 40 dB/Dekade. Zwischen f_{g2} und f_{g3} beträgt die Phase -180° , was eine Umkehr der Funktion bedeutet. An dieser Stelle ist das Ausgangssignal gegenüber dem Eingangssignal invertiert. Ob eine Schaltung bei einer bestimmten Frequenz schwingt, hängt davon ab, ob für die ausgewählte Frequenz die Schwingungsbedingung

$$A_0 = \underline{k} \cdot \underline{A}_D \equiv 1 \Rightarrow \begin{cases} |\underline{A}_0| \equiv |\underline{k}| \cdot |\underline{A}_D| \equiv 1 & \text{Amplitudenbedingung} \\ \varphi \cdot (\underline{k} \cdot \underline{A}_D) \equiv 0^\circ, 360^\circ, \dots & \text{Phasenbedingung} \end{cases} \quad (5.18)$$

mit $\underline{k} = \frac{1}{\underline{A}}$

erfüllt ist. Sie besteht aus den beiden hier aufgeführten Teilen, der Amplitudenbedingung und der Phasenbedingung. Nur wenn beide Bedingungen erfüllt sind, gibt es am Ausgang der Schaltung eine Schwingung mit konstanter Amplitude. Für die Schaltung in Bild 5.12 wäre diese Bedingung für einen invertierenden Verstärker mit $A = 1000$ erfüllt. In Bild 5.13 kann man bei der Frequenz, für die eine Phase von -180° abgelesen werden kann, die Differenzverstärkung $A_D = 10^3$ ermitteln. Eine Invertierung des Eingangssignals erzeugt weitere 180° Phasenverschiebung, so dass beide Bedingungen in Gleichung 5.18 erfüllt sind. Dies gilt aber nur, wenn die zur Gegenkopplung verwendeten Widerstände keine zusätzlichen parasitären Kapazitäten oder Induktivitäten

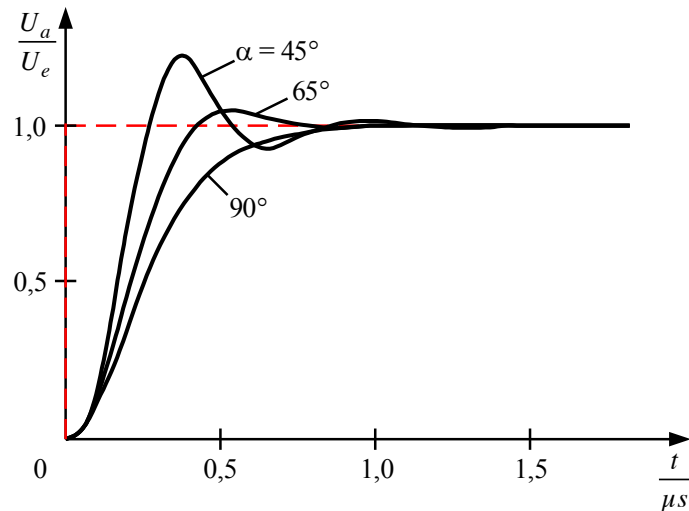


Bild 5.14: Sprungantwort am Ausgang eines Operationsverstärkers in Abhängigkeit der Phasenreserve α . (gestrichelt: U_e)

besitzen. In diesem realen Fall wird die Frequenz, bei der die Schwingungsbedingung erfüllt ist knapp neben der hier ermittelten liegen. Beim Aufbau von Verstärkern muss man also das Bode-Diagramm des verwendeten Operationsverstärkers genau analysieren. Dies kann insbesondere bei breitbandigen Verstärkern Probleme bereiten, da sich die Phase dabei auch mehrfach um 360° verschieben kann, und damit innerhalb des Frequenzbereichs an mehreren Stellen, d.h. bei mehreren Frequenzen die Schwingungsbedingung erfüllt sein kann. Dort muss durch zusätzliche externe Bauelemente eine Korrektur des Frequenz- bzw. Phasengangs vorgenommen werden, so dass die Bedingungen nach Gl. 5.18 im gewünschten Arbeitsbereich nie erfüllt sind. Eine Aussage über die Stabilität einer Verstärkerschaltung liefert die sogenannte Phasenreserve. Diese ist definiert als

$$\alpha = 180^\circ - |\varphi(f)| \quad (5.19)$$

Allgemein lässt sich sagen, je größer die Phasenreserve, umso stabiler verhält sich eine Verstärkerschaltung. Für den Entwickler ist es wichtig, aus den Daten und Diagrammen des Datenblatts die richtigen und notwendigen Informationen zu entnehmen. Aus dem Bode-Diagramm (Bild 5.13) lässt sich für eine gewünschte obere Grenzfrequenz der maximale Verstärkungsfaktor, der durch die externe Gegenkopplung eingestellt werden kann, entnehmen. So ist z.B. für $f_{g3} = 1$ MHz nur noch eine maximale Verstärkung $A < 100$ möglich. Bei den ganzen Betrachtungen sind wir bisher davon ausgegangen, dass rein sinusförmige Signale verstärkt werden. Häufig liegen aber auch trapez- oder rechteckförmige Signale vor. Um alle in diesen Signalen vorkommenden Frequenzen zu ermitteln muss die Fourierreihe der Eingangsfunktion ermittelt werden. Daraus kann dann ein Einschwingvorgang des Ausgangssignals in Abhängigkeit von der Phasenre-

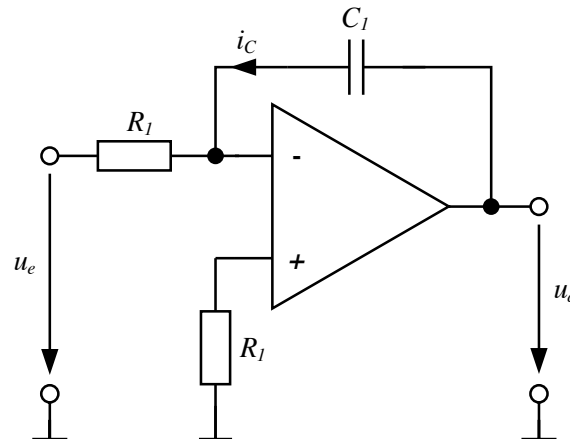


Bild 5.15: Invertierender Integrator mit Eingangsruhestromkompensation.

serve ermittelt werden. Bild 5.14 zeigt die Sprungantwort am Ausgang eines Operationsverstärkers in Abhängigkeit der Phasenreserve α . Bei $\alpha < 90^\circ$ erkennt man ein Überschwingen, das mehr oder weniger schnell abklingt. Man hat in diesem Fall eine Überlagerung des tatsächlichen Signals mit einer gedämpften Schwingung, die durch die Anteile höherer Frequenzen im Eingangssignal entsteht und die Phasenreserve für diese Oberwelle nahe 0° liegt.

5.6 Gegengekoppelte Schaltungen mit Operationsverstärkern

Nachdem die wesentlichen Eigenschaften von Operationsverstärkern erläutert wurden, wollen wir uns jetzt mit einer Reihe von Schaltungen zur Realisierung bestimmter Funktionen mit Operationsverstärkern als zentralem Element befassen. Die Grundschaltungen invertierender und nichtinvertierender Verstärker wurden bereits besprochen. Auf diesen Schaltungen werden die nun folgenden aufbauen.

5.6.1 Invertierender Integrator

Ersetzt man den Rückkopplungswiderstand R_N beim invertierenden Verstärker durch den Kondensator C_1 wird daraus ein invertierender Integrator der in Bild 5.15 dargestellt ist. Der Widerstand am nichtinvertierenden Eingang dient der Kompensation des Eingangsruhestroms. Die Ausgangsspannung u_a ergibt sich zu:

$$u_a = \frac{1}{C_1} \cdot \left[\int_0^t i_C(t) dt + Q_0(t=0) \right] \quad (5.20)$$

Q_0 ist dabei die Ladung, die sich zum Zeitpunkt $t = 0$ auf dem Kondensator befindet. Setzt man für $i_C = -u_e/R_1$, erhält man die Ausgangsspannung u_a in Abhängigkeit von der Eingangsspannung u_e und der Zeitkonstanten $\tau = R_1 \cdot C_1$.

$$u_a = -\frac{1}{R_1 \cdot C_1} \cdot \int_0^t u_e(t) dt + u_a(t=0) \quad (5.21)$$

Betrachten wir nun zwei Sonderbetriebsfälle:

1. Für eine zeitlich konstante Eingangsspannung erhalten wir die Ausgangsspannung

$$u_a = -\frac{u_e}{R_1 \cdot C_1} \cdot t + u_a(t=0) \quad (5.22)$$

Die Ausgangsspannung ändert sich also linear mit der Zeit. Die Schaltung eignet sich deshalb sehr gut zur Erzeugung von Dreieck- und Sägezahnschwingungen.

2. Für eine sinusförmige Eingangsspannung u_e wird die Ausgangsspannung

$$u_a = -\frac{1}{R_1 \cdot C_1} \cdot \int_0^t \hat{U}_e \cdot \sin \omega t dt + u_a(t=0) = \frac{\hat{U}_e}{\omega \cdot R_1 \cdot C_1} \cdot \cos \omega t + u_a(t=0) \quad (5.23)$$

Die Amplitude der Ausgangsspannung ist also umgekehrt proportional zur Kreisfrequenz ω .

5.6.2 Invertierender Differenzierer

Wenn man mit einem Operationsverstärker durch äußere Beschaltung das Eingangssignal integrieren kann, liegt die Vermutung nahe, dass es auch eine Schaltung zur Differenzierung eines Signals geben wird. Vertauscht man Widerstand und Kondensator wie in Bild 5.16 gezeigt, erhält man einen invertierenden Differenzierer. Durch Anwendung der Knotenregel am Summationspunkt erhält man die Ausgangsspannung

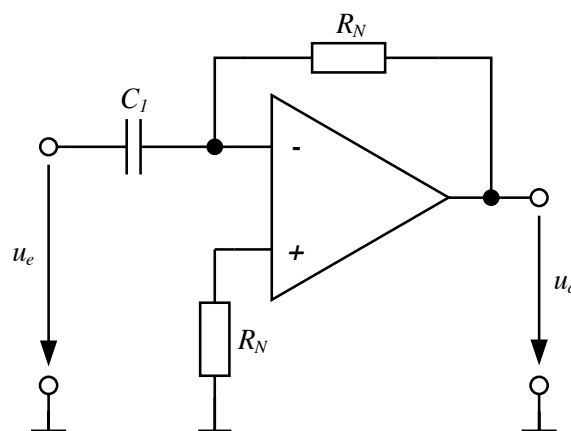


Bild 5.16: Invertierender Differenzierer mit Eingangsruhestromkompensation.

$$\begin{aligned}
C_1 \cdot \frac{\partial u_e}{\partial t} + \frac{u_a}{R_N} &= 0 \\
u_a &= -R_N \cdot C_1 \cdot \frac{\partial u_e}{\partial t}
\end{aligned} \tag{5.24}$$

Auch hier sollen wieder zwei Sonderfälle betrachtet werden:

1. Die Änderung der Eingangsspannung nach der Zeit sei konstant, wie dies bei einer Dreieckschwingung der Fall ist. Die Ausgangsspannung wird damit

$$u_a = -R_N \cdot C_1 \frac{\partial u_e}{\partial t} = -\tau \cdot \frac{\Delta u_e}{\Delta t} = \text{const.} \tag{5.25}$$

2. Im Falle einer sinusförmigen Eingangsspannung wird die Ausgangsspannung

$$u_a = -R_N \cdot C_1 \cdot \frac{\partial(\hat{U}_e \cdot \sin \omega t)}{\partial t} = -\omega \cdot R_N \cdot C_1 \cdot \hat{U}_e \cos \omega t \tag{5.26}$$

5.6.3 Addierer

Häufig ist es notwendig, mehrere Eingangssignale zu einem gemeinsamen Ausgangssignal zusammenzuführen. Dazu kann eine Schaltung nach Bild 5.17 aufgebaut werden. Die Schaltung ist eine Erweiterung des invertierenden Verstärkers um zusätzliche Eingänge. Am invertierenden Eingang ist nach der Knotenregel auch hier die Summe aller Ströme gleich Null. Entsprechend dem invertierenden Verstärker wird die Ausgangsspannung u_a

$$u_a = - \left(\frac{R_N}{R_{11}} \cdot u_{e1} + \frac{R_N}{R_{12}} \cdot u_{e2} + \frac{R_N}{R_{13}} \cdot u_{e3} \right) \tag{5.27}$$

ist jetzt die Summe aller Eingangsspannungen. Wie aus Gleichung 5.27 zu entnehmen ist, kann jede einzelne Eingangsspannung mit einem eigenen Verstärkungsfaktor A gewichtet werden. Auch hier ist es sinnvoll, mit einem Widerstand R_K eine Kompensation des Eingangsruhestroms vorzunehmen. R_K wird ermittelt, in dem man den Widerstandswert der Parallelschaltung aller Widerstände am invertierenden Eingang berechnet.

5.6.4 Subtrahierer

Nicht ganz so einfach wie das Addieren ist das Subtrahieren mehrerer Eingangsspannungen. Zum einen können nur zwei Spannungen voneinander subtrahiert werden, zum andern müssen strenge Entwurfskriterien eingehalten werden. Die Schaltung in Bild 5.18 mit einem einzelnen Operationsverstärker besteht aus einem invertierenden (wenn $u_{e+} = 0$ ist) und einem nichtinvertierenden (wenn $u_{e-} = 0$ ist) Verstärker. Der Spannungsteiler R_3 , R_4 muss die um eins höhere Verstärkung des nichtinvertierenden

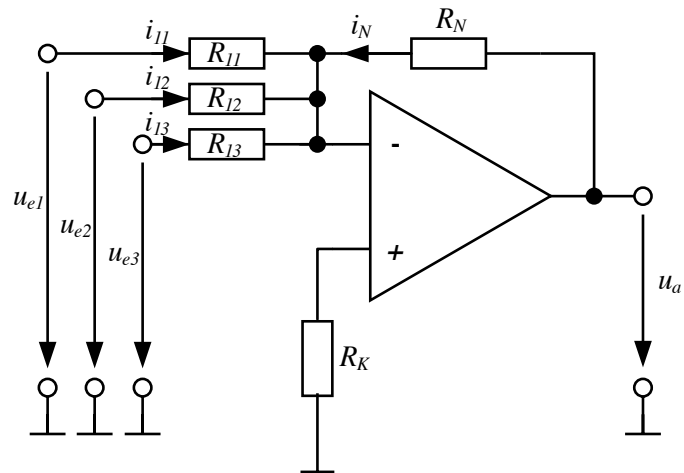


Bild 5.17: Invertierender Addierer mit Eingangsruhestromkompensation.

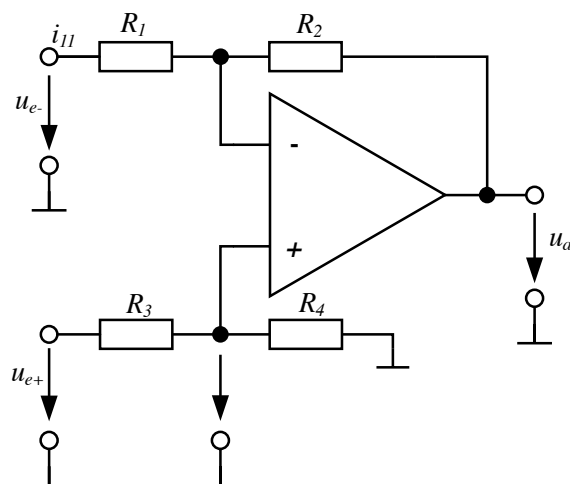


Bild 5.18: Subtrahierer mit einem Operationsverstärker.

Verstärkers kompensieren, damit am Ausgang das richtige Ergebnis erscheint. Es ist also:

$$\begin{aligned}
 u_a &= -\frac{R_2}{R_1} \cdot u_{e-} \quad \text{für } u_{e+} = 0 \quad \text{oder} \\
 u_a &= \left(1 + \frac{R_2}{R_1}\right) \cdot u_+ \quad \text{für } u_{e-} = 0 \quad \text{mit } u_+ = u_{e+} \cdot \frac{R_4}{R_3 + R_4} \\
 &\quad \text{Ist } \frac{R_2}{R_1} = \frac{R_4}{R_3} \quad \text{wird} \\
 u_a &= \frac{A_+}{1 + A_+} \cdot (1 + A_-) \cdot u_+ \quad \text{für } u_{e-} = 0
 \end{aligned} \tag{5.28}$$

Die Ausgangsspannung ist damit:

$$u_a = \left(1 + \frac{R_2}{R_1}\right) \cdot u_+ - \frac{R_2}{R_1} \cdot u_{e-} = A \cdot (u_{e+} - u_{e-}) \quad \text{mit } A = \frac{R_2}{R_1} \tag{5.29}$$

Sind die beiden Verstärkungsfaktoren A_+ und A_- ungleich groß, so ist die Ausgangsspannung nicht mehr exakt die Differenz der beiden Eingangsspannungen. Sie wird zu

$$u_a = \frac{1 + A_-}{1 + A_+} \cdot A_+ \cdot u_{e+} - A_- \cdot u_{e-} \tag{5.30}$$

Für einen optimalen Betrieb der Schaltung ist es also wichtig, dass die Widerstandsverhältnisse in den beiden Zweigen sehr genau eingehalten werden und beide Verstärkungsfaktoren gleich groß sind.

5.6.5 Logarithmierer

Soll die Ausgangsspannung eines Verstärkers dem Logarithmus der Eingangsspannung entsprechen, kann man die Exponentialfunktion einer Diode nutzen. Da der Diodenstrom im Durchlassbereich aber den Korrekturfaktor m ($1 < m < 2$) besitzt, ist diese Exponentialfunktion nur eingeschränkt tauglich. Beim Kollektorstrom eines Bipolartransistors jedoch ist dieser Korrekturfaktor nicht mehr vorhanden. Die Schaltung in Bild 5.19 ist das Ergebnis dieser Vorbetrachtung. Für den Kollektorstrom I_C gilt:

$$I_C = I_{CS} \cdot \left(e^{\frac{U_{BE}}{U_T}} - 1 \right) \approx I_{CS} \cdot e^{\frac{U_{BE}}{U_T}} \Rightarrow U_{BE} = U_T \cdot \ln \frac{I_C}{I_{CS}} \tag{5.31}$$

Der Kollektorstrom wird bestimmt durch die Eingangsspannung u_e und den Widerstand R_1 , da wir annehmen können, dass der Eingangsstrom des Operationsverstärkers vernachlässigbar klein ist. Damit ergibt sich die Ausgangsspannung zu:

$$u_a = -U_T \cdot \ln \frac{u_e}{I_{CS} \cdot R_1} = -U_T \cdot \ln 10 \cdot \log \frac{u_e}{I_{CS} \cdot R_1} \approx -60 \text{ mV} \cdot \log \frac{u_e}{I_{CS} \cdot R_1} \tag{5.32}$$

für Eingangsspannungen $u_e > 0$. Durch den hohen Eingangswiderstand des Operations-

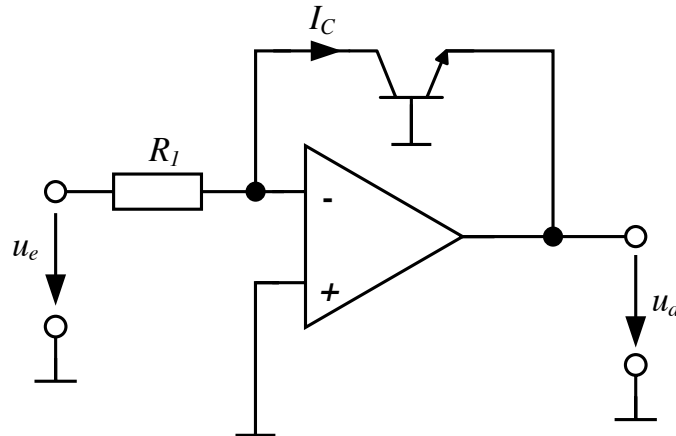


Bild 5.19: Grundsaltung eines Logarithmierers.

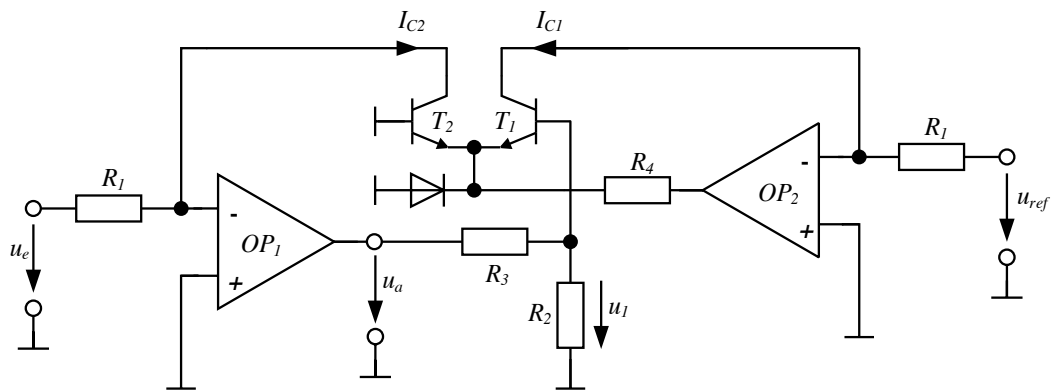


Bild 5.20: Temperaturkompensierter Logarithmierer.

verstärkers ergibt sich ein Gesamtbereich des Logarithmus von bis zu 9 Dekaden. Die Grundsaltung in Bild 5.19 hat aber den Nachteil, dass die Temperaturabhängigkeit des Transistors voll in die Ausgangsspannung eingeht. Soll dies vermieden werden, ist eine erheblich aufwendigere Schaltung notwendig. Bild 5.20 zeigt die vollständige Schaltung eines temperaturkompensierten Logarithmierers mit zwei Operationsverstärkern. Hier wird die Differenz zweier Logarithmen gebildet, wodurch sich die Temperatureinflüsse eliminieren lassen. Der Differenzverstärker aus den beiden Transistoren T_1 und T_2 bildet nun die eigentliche Logarithmierung. Da die Wirkungsweise der Schaltung nicht offensichtlich ist, soll diese nun etwas genauer untersucht werden. Aus der Maschenregel folgt für die Spannungen am Differenzverstärker

$$u_1 + U_{BE2} - U_{BE1} = 0 \quad (5.33)$$

Für die Kollektorströme der beiden Transistoren gilt:

$$I_{C1} = I_{CS} \cdot e^{\frac{U_{BE1}}{U_T}} \quad \text{und} \quad I_{C2} = I_{CS} \cdot e^{\frac{U_{BE2}}{U_T}} \quad (5.34)$$

Das Verhältnis der beiden Ströme ist

$$\frac{I_{C1}}{I_{C2}} = \frac{I_{CS} \cdot e^{\frac{U_{BE1}}{U_T}}}{I_{CS} \cdot e^{\frac{U_{BE2}}{U_T}}} = e^{\frac{U_{BE1} - U_{BE2}}{U_T}} = e^{\frac{u_1}{U_T}} \quad (5.35)$$

Aus Bild 5.20 können wir für die beiden Kollektorströme und die Spannung u_1 folgende Beziehungen entnehmen:

$$I_{C1} = \frac{u_{ref}}{R_1} \quad , \quad I_{C2} = \frac{u_e}{R_1} \quad , \quad u_1 = \frac{R_2}{R_2 + R_3} \cdot u_a \quad (5.36)$$

Durch Einsetzen erhält man die Ausgangsspannung

$$u_a = -U_T \cdot \frac{R_2 + R_3}{R_2} \cdot \ln \frac{u_e}{u_{ref}} \quad (5.37)$$

Der Widerstand R_4 hat keinen Einfluss auf die Ausgangsspannung. Er sollte so gewählt werden, dass der Spannungsabfall kleiner bleibt, als die maximale Ausgangsspannung des Operationsverstärkers OP2. Der Widerstand R_2 sollte ebenfalls nicht zu hoch ohmig sein. Werte um 1 k Ω sind hierbei üblich. Wenn man auch noch den Einfluss von U_T reduzieren möchte, muss R_2 einen positiven Temperaturkoeffizienten von etwa 0,3 %/K besitzen.

5.6.6 Exponentialfunktion

Eine Schaltung zur Erzeugung einer e -Funktion ist ähnlich aufgebaut wie der gerade besprochene Logarithmierer. Wie bei Integrator und Differenzierer werden einfach die Bauelemente getauscht. Bild 5.21 zeigt einen einfachen e -Funktionsgenerator.

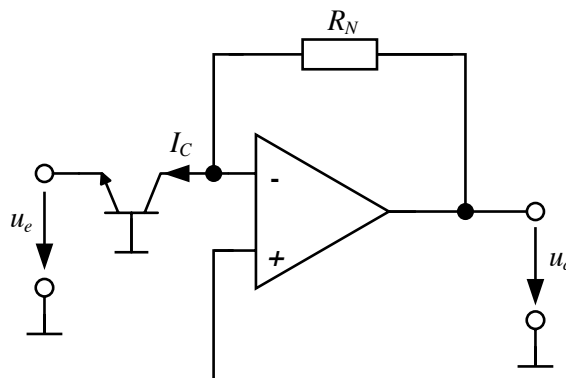


Bild 5.21: Einfacher e -Funktionsgenerator.

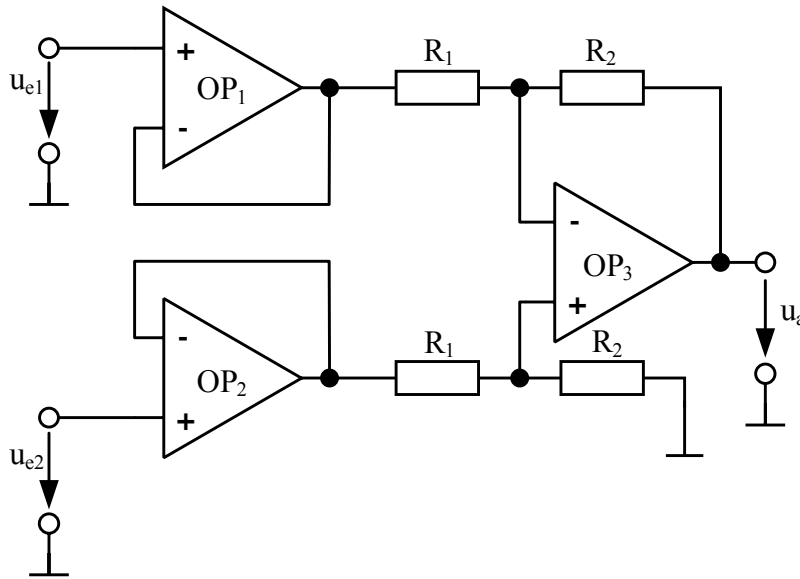


Bild 5.22: Subtrahierer mit vorgeschalteten Spannungsfolgern.

Die Ausgangsspannung der Schaltung wird mit

$$I_C = I_{CS} \cdot e^{\frac{U_{BE}}{U_T}} = I_{CS} \cdot e^{\frac{-u_e}{U_T}} \quad (5.38)$$

zu

$$u_a = I_C \cdot R_N = R_N \cdot I_{CS} \cdot e^{\frac{-u_e}{U_T}} \quad (5.39)$$

Die Eingangsspannung u_e muss hier negativ sein. Für die Temperaturabhängigkeit der Schaltung gilt die gleiche Aussage wie beim Logarithmierer. Die Schaltung eines temperaturkompensierten e-Funktionsgenerators ist genau so aufwendig, wie die des Logarithmierers in Bild 5.20.

5.6.7 Messverstärker

Zur genauen Messung kleiner, störbehafteter Signale ist es nicht ausreichend, Schaltungen mit nur einem Operationsverstärker, wie sie bisher besprochen wurden, einzusetzen. Zwar haben Operationsverstärker an sich eine recht gute Gleichtaktunterdrückung, aber die genaue Messung von Potentialdifferenzen ist mit einem reinen Subtrahierer nicht immer möglich, da der Eingangswiderstand im Wesentlichen durch die externe Beschaltung bestimmt wird und recht niedrig (im Verhältnis zum Eingangswiderstand des OPV) sein kann. Deshalb sollten wenigstens Spannungsfolger als Impedanzwandler vorgeschaltet werden, wie in Bild 5.22 dargestellt. Die Ausgangsspannung ist:

$$u_a = \frac{R_2}{R_1} \cdot (u_{e2} - u_{e1}) \quad (5.40)$$

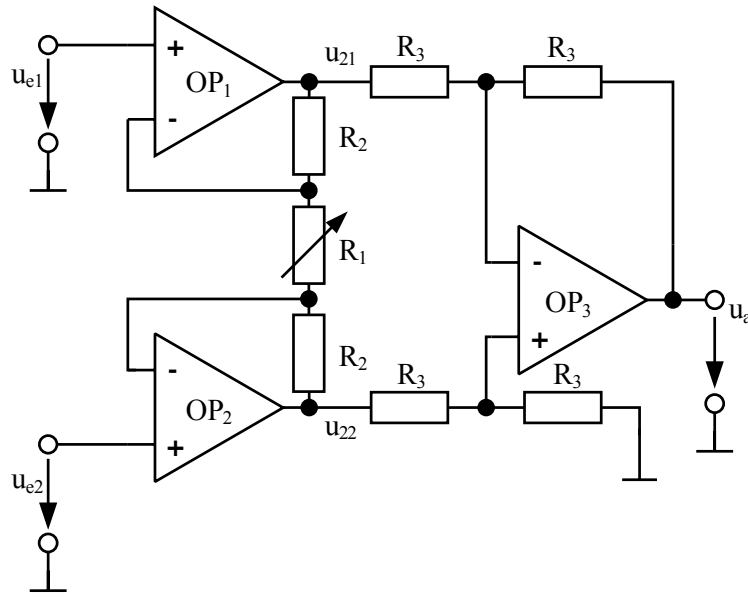


Bild 5.23: Schaltung eines Instrumentation Amplifiers (Elektrometersubtrahierer).

Eine höhere Gleichtaktunterdrückung erhält man, wenn man anstelle der Spannungsfolger nichtinvertierende Verstärker, die einen sehr hohen Eingangswiderstand besitzen, einsetzt. Die zugehörige Schaltung ist in Bild 5.23 gezeigt. Am Widerstand R_1 tritt die Potentialdifferenz $u_{e2} - u_{e1}$ auf. Das Verhältnis der Widerstände R_2/R_1 bestimmt die Verstärkung der Schaltung, da der nachgeschaltete Subtrahierer nur die Verstärkung $A = 1$ aufweist. Diese Schaltung wird als „Instrumentation Amplifier“ oder Elektrometersubtrahierer bezeichnet. Die Ausgangsspannung dieser Schaltung ist:

$$u_a = u_{22} - u_{12} = \left(1 + \frac{2 \cdot R_2}{R_1}\right) \cdot (u_{e2} - u_{e1}) \quad (5.41)$$

Diese Schaltung kann zwar mit diskreten Bauelementen aufgebaut werden, eine integrierte Version, die nur noch den Widerstand R_1 als externes Bauelement zulässt ist aber erheblich genauer. Damit ist es möglich, mit nur einem Widerstand die Verstärkung einzustellen. Bei reiner Gleichtaktansteuerung d.h. $u_{e1} = u_{e2} = u_G$, wird auch $u_{21} = u_{22} = u_G$. Damit werden Gleichtaktsignale mit $A_G = 1$ verstärkt, unabhängig von der Verstärkung A_D der Differenzsignale. Die Gleichtaktunterdrückung der Gesamtschaltung ist von der letzten Stufe geprägt. Je genauer die Widerstandspaare R_3 im invertierenden und im nichtinvertierenden Zweig angepasst sind, um so besser wird die Gleichtaktunterdrückung. Sie kann um einige Größenordnungen höher sein, als die eines einzelnen Operationsverstärkers.

6 Kippschaltungen

Lernziele

- Kennenlernen des prinzipiellen Aufbaus von Kippschaltungen
- Funktion von Kippschaltungen mit bipolaren Transistoren
- Funktion von Kippschaltungen mit Operationsverstärkern

6.1 Kippschaltungen mit bipolaren Transistoren

Bei den bisher besprochenen linearen Schaltungen wurde der Arbeitspunkt so ausgewählt, dass eine kleine Aussteuerung um den Arbeitspunkt zu einer proportional verstärkten Ausgangsspannung geführt hat. Bei den Kippschaltungen arbeiten wir mit zwei Betriebszuständen der Transistoren. Dabei basieren die Kippschaltungen auf einer Mitkopplung. Sie unterscheiden sich von mitgekoppelten Linearschaltungen (Oszillatoren) dadurch, dass sich ihre Ausgangsspannung nicht kontinuierlich ändert, sondern nur zwischen zwei festen Spannungswerten hin und her springt. Da die Ausgangsspannung nur diese beiden Werte annehmen kann, werden die Kippschaltungen oft den digitalen Schaltungen zugeordnet. Sie bilden den Grenzbereich zwischen den analogen und digitalen Schaltungen. Der Kippvorgang kann auf unterschiedliche Weise ausgelöst werden. Bei den bistabilen Kippschaltungen ändert sich der Ausgangszustand nur dann, wenn mit Hilfe des Eingangssignals ein Umkippvorgang ausgelöst wird. Beim Flip-Flop genügt ein kurzer Impuls, während beim Schmitt-Trigger ein beständiges Eingangssignal erforderlich ist. Eine monostabile Kippschaltung besitzt nur einen stabilen Zustand. Der zweite Zustand ist nur für eine bestimmte Zeit, die durch die Dimensionierung der Bauelemente festgelegt wird, stabil. Nach Ablauf dieser Zeit kippt die Schaltung wieder zurück in ihre Ausgangszustand. Eine astabile Kippschaltung besitzt keinen stabilen Zustand, sondern kippt ohne äußere Anregung ständig hin und her. Sie wird auch als Multivibrator bezeichnet. Die drei Grundtypen von Kippschaltungen können durch unterschiedliche Auswahl der Mitkopplungsglieder realisiert werden. Eine Übersicht über die verschiedenen Koppelglieder ist in Tabelle 6.1 gegeben. Bild 6.1 zeigt den schematischen Aufbau einer Kippschaltung.

Kippschaltung	Name	Koppelglied 1	Koppelglied 2
Bistabil	Flip-Flop Schmitt-Trigger	R	R
Monostabil	Univibrator	R	C
Astabil	Multivibrator	C	C

Tabelle 6.1: Realisierung der Koppelglieder für verschiedene Kippschaltungen.

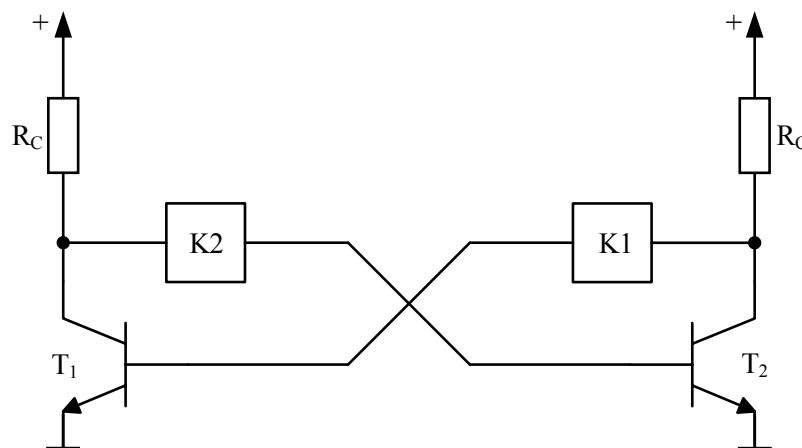


Bild 6.1: Prinzipielle Anordnung von Kippschaltungen.

6.1.1 Bistabile Kippschaltungen

6.1.1.1 Flip-Flop

Die Funktionsweise des Flip-Flop nach Bild 6.2 kann wie folgt beschrieben werden: Eine positive Spannung am Eingang U_{e1} öffnet den Transistor T_1 , damit fließt ein Kollektorstrom durch T_1 und die Kollektorspannung an T_1 wird kleiner. Damit wird die Basis-Emitterspannung an T_2 kleiner, d.h. der Basisstrom von T_2 sinkt. Damit reduziert sich der Kollektorstrom von T_2 und somit nimmt die Kollektorspannung von T_2 zu. Über den Widerstand R_{B1} erfolgt eine Mitkopplung, was zu einem Anstieg des Basisstroms von T_1 führt. Der stationäre Zustand ist erreicht, wenn die Kollektorspannung von T_1 bis auf die Sättigungsspannung abgenommen hat. Damit sperrt T_2 und T_1 wird über den Widerstand R_{B1} leitend gehalten. Deshalb kann man nach Abschluss des Umkippvorganges die Spannung am Eingang U_1 wieder auf Null setzen, ohne dass die Schaltung ihren Zustand ändern wird. Man kann das Flip-Flop wieder zurückkippen, indem man eine positive Spannung an den Eingang U_{e2} anlegt. Dann wird sich der oben beschriebene Umkippvorgang in entgegengesetzter Abfolge wiederholen bis entsprechend der Transistor T_2 leitet und T_1 sperrt. In der Digitaltechnik wird das Flip-Flop zur Speicherung von Zuständen (Informationen) benutzt. Die Eingänge U_{e1} und U_{e2} werden zum

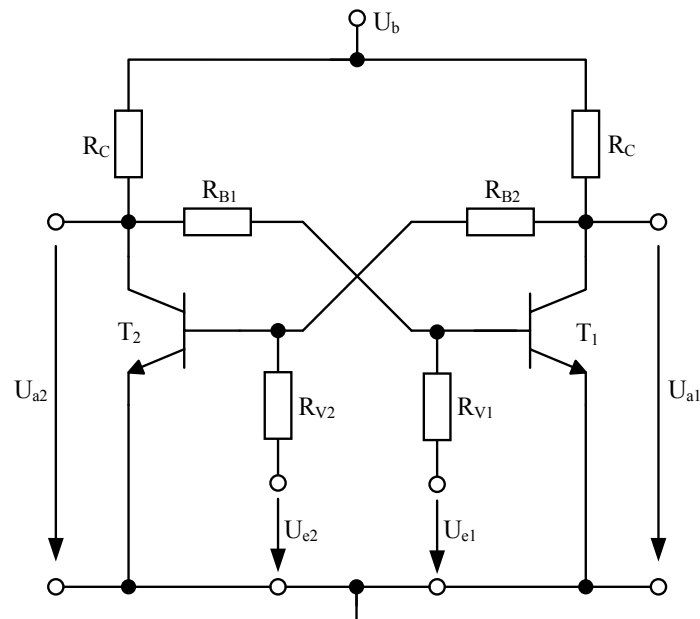


Bild 6.2: Schaltung eines Flip-Flop.

„Setzen“ bzw. „Rücksetzen“ benutzt und deshalb entsprechend als S- bzw. R-Eingänge bezeichnet. Das Flip-Flop wird dann RS-Flip-Flop genannt.

6.1.1.2 Schmitt-Trigger

Eine andere Möglichkeit, einen Kippvorgang durchzuführen, besteht darin, nur eine Eingangsspannung zu verwenden. Dabei leitet man den Kippvorgang dadurch ein, dass man die Eingangsspannung abwechselnd positiv bzw. negativ macht. Bild 6.3 zeigt die Schaltung eines Schmitt-Triggers und die Abhängigkeit der Ausgangsspannung von der Eingangsspannung. Wenn die Eingangsspannung die obere Schaltschwelle $U_{e, \text{ein}}$ überschreitet, springt die Ausgangsspannung an den oberen Grenzwert $U_{a, \text{max}}$. Sie springt erst wieder auf null zurück, wenn die untere Schaltschwelle $U_{e, \text{aus}}$ erreicht ist. Dieser Effekt ist in der Übertragungskennlinie (Bild 6.3b) als Hysterese zu sehen. Die Schaltung funktioniert wie nachfolgend beschrieben: Durch eine Ansteuerung des Transistors T_2 der bistabilen Kippschaltung mit einer Spannung U_e über einen Basiswiderstand R_B erhält man einen zweistufigen Transistorverstärker mit einer Mitkopplung nach Bild 6.3a. Für $U_e = 0 \text{ V}$ gibt es zwei stabile Zustände. Erreicht die Eingangsspannung die Schwelle $U_{e, \text{ein}}$, wird die Ausgangsspannung U_a infolge der Mitkopplung über R_2 sprunghaft den Wert $U_{a, \text{max}}$ erreichen, auch wenn sich die Eingangsspannung nur langsam ändert. U_a springt erst dann wieder auf den Wert $U_{a, \text{min}}$ zurück, wenn die Eingangsspannung die Schwelle $U_{e, \text{aus}}$ erreicht. Die Spannungsdifferenz $U_{e, \text{ein}} - U_{e, \text{aus}}$ nennt man auch Schalthysterese, die umso kleiner wird, je größer die Abschwächung des mitgekoppelten Ausgangssignals durch den Spannungsteiler R_2, R_B ist. Die Schalthysterese verschwindet für $R_B = 0$. Der Schmitt-Trigger kann also sinusförmige oder dreieckförmige Eingangs-

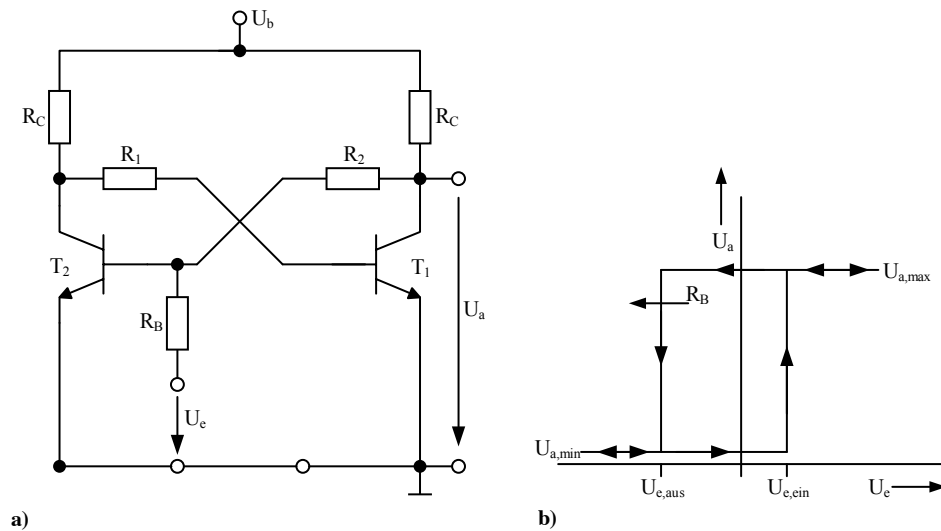


Bild 6.3: Schaltung eines Schmitt-Triggers (a) und Übertragungskennlinie (b).

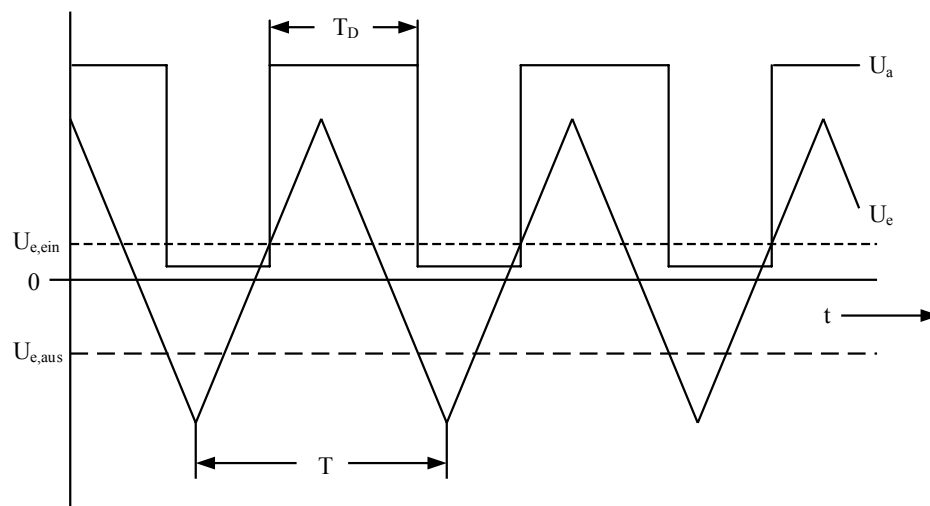


Bild 6.4: Umwandlung einer dreieckförmigen Eingangsspannung in eine rechteckförmige Ausgangsspannung.

spannungen in Rechtecksignale nach Bild 6.4 umwandeln.

Während bei der Schaltung nach Bild 6.3 eine negative Eingangsspannung für das Erreichen der Ausschaltschwelle benötigt wird, ist ein emittergekoppelter Schmitt-Trigger mit ausschließlich positiven Spannungen ansteuerbar und somit sehr einfach in der Digitaltechnik einsetzbar. Deshalb soll die Funktion der Schaltung an dieser Stelle ausführlich diskutiert werden. Eine entsprechende Schaltung ist in Bild 6.5 zusammen mit der Übergangskennlinie dargestellt. Geht man von der Annahme aus, dass $R_{B2} \gg R_{C2}$ ist und die Restspannung eines eingeschalteten Transistors $U_{CE} \approx 0$ ist, kann man folgende vereinfachten Überlegungen anstellen: Ist die Eingangsspannung $U_e < U_{e, \text{ein}}$, so sperrt Transistor T_1 und Transistor T_2 leitet. Damit ist $I_{E1} \ll I_{E2}$. Die Spannung am Widerstand R_E ist

$$U_E = (I_{E1} + I_{E2}) \cdot R_E \quad (6.1)$$

oder vereinfacht

$$U_E \approx I_{E2} \cdot R_E. \quad (6.2)$$

Mit

$$I_{E2} \approx \frac{U_b}{R_E + R_{C2}} \quad (6.3)$$

wird die Spannung für die Einschaltsschwelle zu

$$U_{e, \text{ein}} = U_b \cdot \frac{R_E}{R_E + R_{C2}} + U_{BE, T_1} \quad (6.4)$$

Für $U_e > U_{e, \text{ein}}$ leitet Transistor T_1 während Transistor T_2 sperrt. Der Strom I_{E2} ist jetzt sehr klein gegenüber dem Strom I_{E1} in Transistor T_1 .

$$I_{E1} \approx \frac{U_b}{R_E + R_{C1}} \quad (6.5)$$

und damit die Spannung

$$U_E \approx I_{E1} \cdot R_E \quad (6.6)$$

Ist $R_{C1} > R_{C2}$, wird $I_{E1} < I_{E2}$ und damit wird auch die Spannung der unteren Schaltschwelle $U_{e, \text{aus}} < U_{e, \text{ein}}$.

$$U_{e, \text{aus}} = U_b \cdot \frac{R_E}{R_E + R_{C1}} + U_{BE, T_2}. \quad (6.7)$$

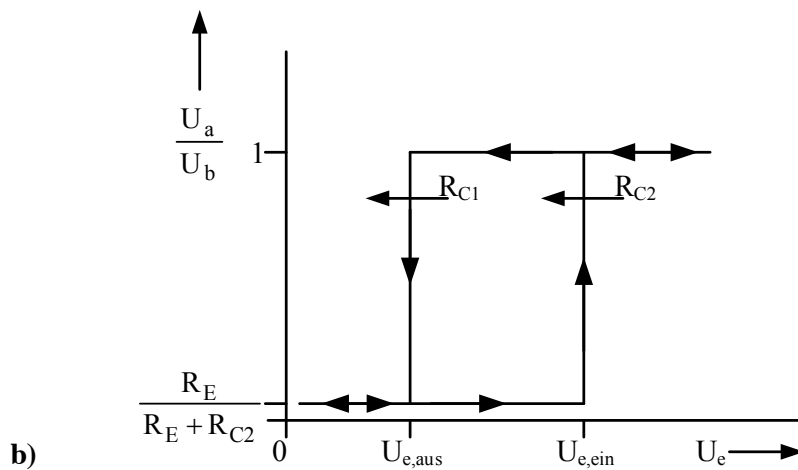
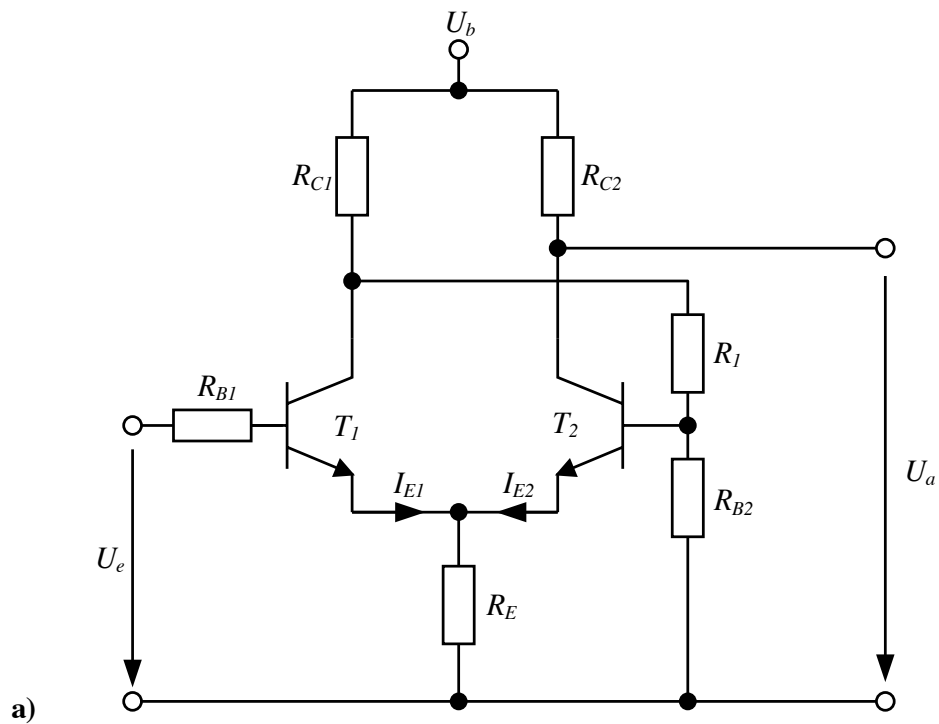


Bild 6.5: Emittergekoppelter, nichtinvertierender Schmitt-Trigger: a) Schaltbild und b) Übertragungskennlinie.

Für die Ansteuerung von Transistor T_2 muss dabei die folgende Bedingung erfüllt sein:

$$U_{B,T_2} = U_b \cdot \frac{R_{B2}}{R_{B2} + R_1 + R_{C1}} = I_{E2} \cdot R_E + U_{BE,T_2} \quad (6.8)$$

6.1.1.3 Monostabile Kippschaltungen

Zur Realisierung eines Univibrators geht man von einem Flip-Flop aus und ersetzt einen Mitkopplungswiderstand durch einen Kondensator. Da über ihn kein Gleichstrom fließen kann, ist T_1 leitend und T_2 sperrt. Bild 6.6 zeigt die Schaltung und den zeitlichen Verlauf der Spannung an verschiedenen Knoten der Schaltung. Ein positiver Spannungsimpuls am Eingang U_e öffnet T_2 . Dadurch springt die Kollektorspannung von T_2 auf null. Dieser Sprung wird über das Hochpassglied $R_{B1} \cdot C_1$ auf die Basis von T_1 übertragen und T_1 sperrt. Über den Mitkopplungswiderstand R_2 wird T_2 leitend gehalten, auch wenn die Eingangsspannung wieder null ist. Über den Widerstand R_{B1} wird der Kondensator C_1 wieder aufgeladen und damit steigt die Basisspannung an T_1 nach folgender Formel an:

$$U_{B1}(t) = U_b \cdot \left(1 - 2 \cdot e^{-\frac{t}{R_{B1} \cdot C_1}}\right) \quad (6.9)$$

Die erforderliche Zeit, damit U_{B1} null wird beträgt:

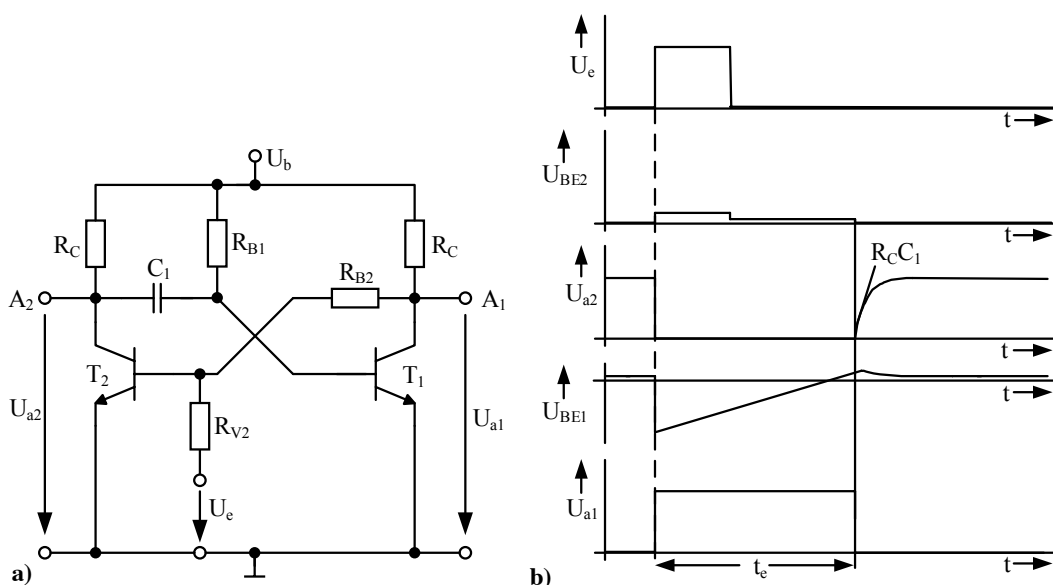


Bild 6.6: Schaltung eines Univibrators (a) und zeitlicher Verlauf der Spannungen an ausgewählten Knoten der Schaltung (b).

$$t_e \approx R_{B1} \cdot C_1 \cdot \ln 2 \approx 0,7 \cdot R_{B1} \cdot C_1 \quad (6.10)$$

Nach dieser Zeit wird der Transistor T_1 wieder leitend, d.h. die Schaltung kippt wieder in ihren stabilen Zustand zurück.

6.1.1.4 Astabile Kippschaltungen

Ersetzt man beim Univibrator (monostabile Kippschaltung) auch den zweiten Rückkoppelwiderstand durch einen Kondensator, werden beide Zustände nur für eine begrenzte Zeit stabil. Die Schaltung kippt zwischen den beiden Zuständen hin und her und wird deshalb Multivibrator genannt. Bild 6.7 zeigt die Schaltung eines Multivibrators und den zeitlichen Verlauf der Spannungen an ausgewählten Knoten der Schaltung. Die Schaltzeiten der beiden Zustände ergeben sich entsprechend Gl. 6.11 und 6.12:

$$t_1 \approx R_{B1} \cdot C_1 \cdot \ln 2 \approx 0,7 \cdot R_{B1} \cdot C_1 \quad (6.11)$$

$$t_2 \approx R_{B2} \cdot C_2 \cdot \ln 2 \approx 0,7 \cdot R_{B2} \cdot C_2 \quad (6.12)$$

Anhand des zeitlichen Verlaufes ist zu sehen, dass die Schaltung kippt, wenn der bisher gesperrte Transistor leitend wird. Bei der Dimensionierung der Widerstände R_{B1} und R_{B2} ist wenig Freiraum. Sie müssen niederohmig gegenüber $\beta \cdot R_C$ sein, damit der Strom ausreichend ist, um den leitenden Transistor in die Sättigung zu bringen. Andererseits müssen sie hochohmig gegenüber R_C sein, damit sich die Kondensatoren bis auf die Betriebsspannung aufladen können. Damit ergibt sich die Beziehung:

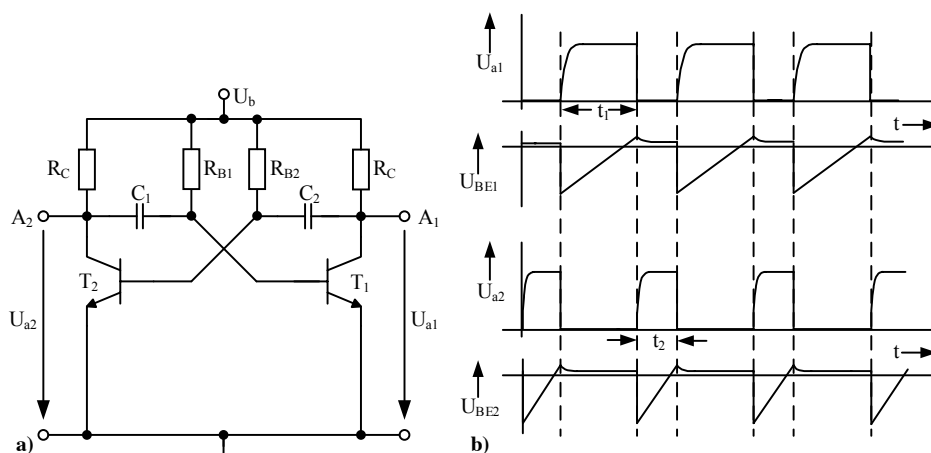


Bild 6.7: Schaltung eines Multivibrators (a) und zeitlicher Verlauf der Spannungen an ausgewählten Knoten der Schaltung (b).

$$R_C \ll R_{B1} \quad \text{und} \quad R_{B2} \ll \beta \cdot R_C \quad (6.13)$$

6.1.2 Kippschaltungen mit Komparatoren

Betrieibt man einen Operationsverstärker ohne Gegenkopplung, wie in Bild 6.8 dargestellt, erhält man einen Komparator. Die Eingangsspannung U_e wird mit einer Referenzspannung U_{ref} verglichen. Die Spannung $U_D = U_e - U_{ref}$ wird mit der Leerlaufverstärkung A_D verstärkt. Die Ausgangsspannung wird:

$$\begin{aligned} U_a &= U_{a,max} & \text{für} & & U_D = U_e - U_{ref} > 0 & \text{und} \\ U_a &= U_{a,min} & \text{für} & & U_D = U_e - U_{ref} < 0 \end{aligned} \quad (6.14)$$

Wegen der sehr hohen Verstärkung eines Operationsverstärkers spricht die Schaltung bereits auf sehr kleine Spannungsdifferenzen an. Sie eignet sich daher zum Vergleich von zwei Spannungen mit hoher Präzision.

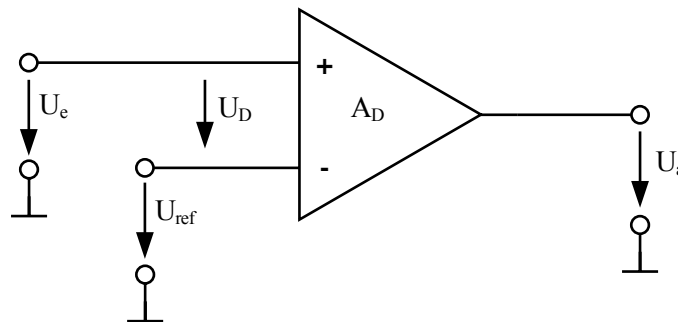


Bild 6.8: Schaltung eines Komparators mit Operationsverstärker.

6.2 Kippschaltungen mit Komparatoren

Im Kapitel 6.1.1 haben wir gezeigt, dass ein Schmitt-Trigger ein Komparator ist, bei dem Ein- und Ausgangsschaltswelle nicht zusammenfallen.

6.2.1 Invertierender Schmitt-Trigger

Den schaltungstechnischen Aufbau eines invertierenden Schmitt-Triggers mit einem Komparator und zwei Widerständen zeigt Bild 6.9. Die Schalthysterese ΔU_e wird durch eine Mitkopplung des Ausgangssignals auf den nichtinvertierenden Eingang des Komparators über den Spannungsteiler aus R_1 und R_2 erreicht. Die beiden Schaltschwellen $U_{e, ein}$ und $U_{e, aus}$ sind bei dieser Schaltung

$$U_{e, ein} = \frac{R_1}{R_1 + R_2} \cdot U_{a, min} \quad , \quad U_{e, aus} = \frac{R_1}{R_1 + R_2} \cdot U_{a, max} \quad (6.15)$$

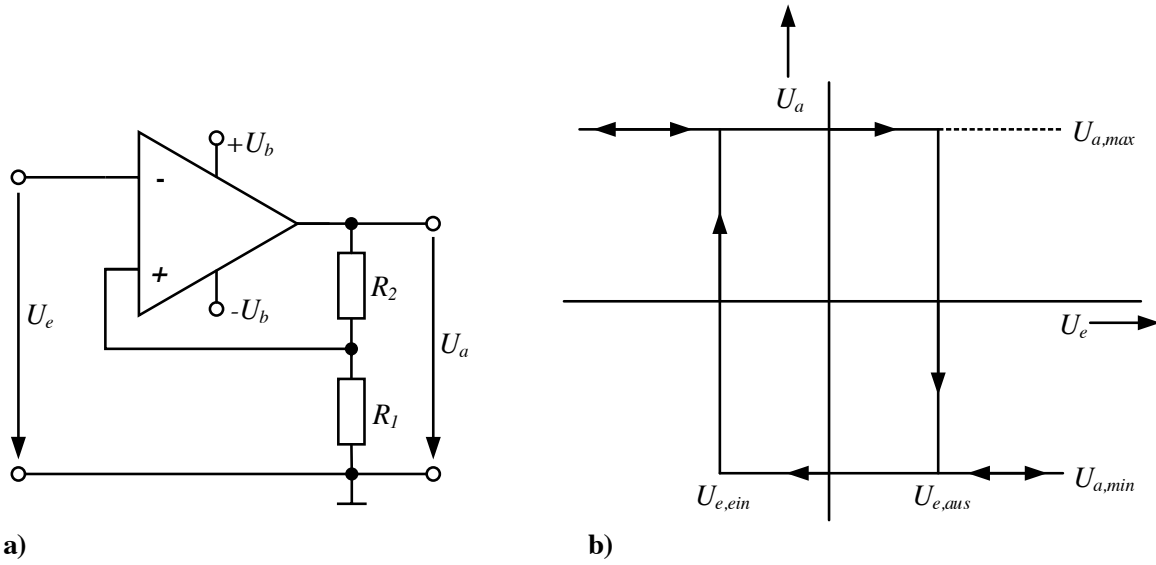


Bild 6.9: Invertierender Schmitt-Trigger mit einem Komparator: a) Schaltbild und b) Übertragungskennlinie.

Daraus ergibt sich der Spannungshub oder die Schalthysterese ΔU_e zu

$$\Delta U_e = \frac{R_1}{R_1 + R_2} \cdot (U_{a,max} - U_{a,min}) \quad (6.16)$$

6.2.2 Nichtinvertierender Schmitt-Trigger

Einen nichtinvertierenden Schmitt-Trigger kann man z.B. mit einem Komparator, zwei Widerständen und einer Referenzspannungsquelle U_{ref} nach Bild 6.10a realisieren. Die Schaltschwellen sind abhängig vom Verhältnis der beiden Widerstände R_1 und R_2 und den maximal erreichbaren Spannungspegeln der Ausgangsspannung U_a (Bild 6.10b). Die Referenzspannungsquelle erlaubt eine Verschiebung der Schaltschwellen. Die Einschaltsschwelle für die Schaltung in Bild 6.10a ergibt sich für $U_D = 0$ aus:

$$\frac{U_{ref} - U_{a,min}}{R_2} = \frac{U_{e,ein} - U_{ref}}{R_1} \quad (6.17)$$

Damit wird:

$$U_{e,ein} = -\frac{R_1}{R_2} \cdot U_{a,min} + U_{ref} \cdot \left(\frac{R_1}{R_2} + 1 \right) \quad (6.18)$$

Entsprechend erhält man für die Ausschaltsschwelle:

$$\frac{U_{a,max} - U_{ref}}{R_2} = \frac{U_{ref} - U_{e,aus}}{R_1} \quad (6.19)$$

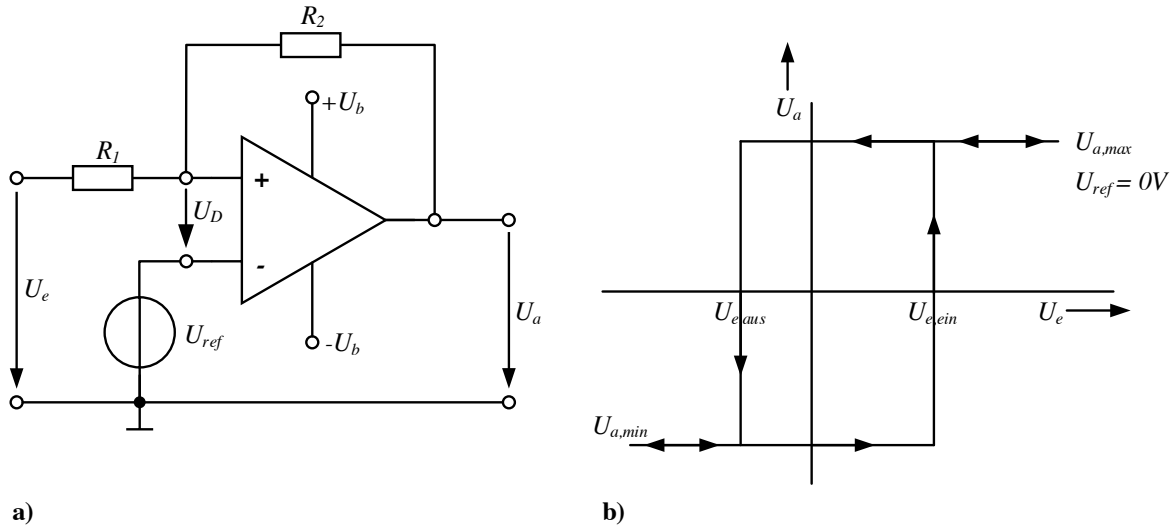


Bild 6.10: Nichtinvertierender Schmitt-Trigger mit einem Komparator: a) Schaltbild und b) Übertragungskennlinie.

und daraus

$$U_{e,aus} = -\frac{R_1}{R_2} \cdot U_{a,max} + U_{ref} \cdot \left(\frac{R_1}{R_2} + 1 \right) \quad (6.20)$$

Aus den Gleichungen 6.18 und 6.20 erkennt man, dass die Schaltschwellen mit U_{ref} sowohl in positiver, als auch in negativer Richtung mit gleicher Gewichtung „verschoben“ werden können. Verbindet man den invertierenden Eingang des Komparators direkt mit Masse, entspricht dies einer Referenzspannung $U_{ref} = 0$ und man erhält eine um den Nullpunkt symmetrische Hysterese. Der Spannungshub ΔU_e der Hysterese ist

$$\Delta U_e = U_{e,ein} - U_{e,aus} = \frac{R_1}{R_2} \cdot (U_{a,max} - U_{a,min}) \quad (6.21)$$

und damit unabhängig von der Referenzspannung.

6.3 Multivibratoren mit Operationsverstärker

Beschaltet man einen Schmitt-Trigger so, dass das Ausgangssignal verzögert auf den Eingang gelangt, entsteht ein Multivibrator entsprechend Bild 6.11a. Wenn die Spannung am invertierenden Eingang den Triggerpegel überschreitet, kippt die Schaltung um. Damit steigt die Ausgangsspannung auf die entgegengesetzte Aussteuerungsgrenze. In Folge läuft die Spannung am invertierenden Eingang in die entgegengesetzte Richtung, bis der andere Schwellwert erreicht ist. Dann kippt die Schaltung in den Ausgangszustand zurück. Der Spannungsverlauf in Bild 6.11b verdeutlicht diesen Vorgang.

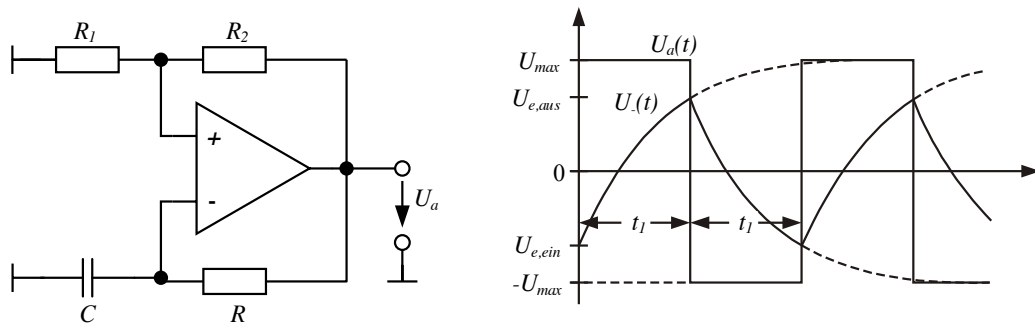


Bild 6.11: Multivibrator mit Komparator: a) Schaltbild, b) zeitlicher Spannungsverlauf.

$$U_{e,ein} = -\alpha \cdot U_{max}$$

$$U_{e,aus} = \alpha \cdot U_{max}$$

$$\text{mit } \alpha = \frac{R_1}{R_1 + R_2}$$

Aus der Schaltung können wir die Differentialgleichung für die Spannung am invertierenden Eingang, U_- aufschreiben:

$$\frac{\partial U_-}{\partial t} = \frac{\pm U_{max} - U_-}{R \cdot C} \quad (6.22)$$

Mit der Randbedingung:

$$U_{e,ein} = -\alpha \cdot U_{max} \quad (6.23)$$

erhält man die folgende Lösung:

$$U_-(t) = U_{max} \cdot \left(1 - (1 + \alpha) \cdot e^{-\frac{t}{R \cdot C}} \right) \quad (6.24)$$

Mit dieser Gleichung lässt sich die Zeit berechnen, die zum Erreichen des Triggerpegels $U_{e,aus} = \alpha \cdot U_{max}$ benötigt wird.

$$t_1 = R \cdot C \cdot \ln \frac{1 + \alpha}{1 - \alpha} = R \cdot C \cdot \ln \left(1 + \frac{2 \cdot R_1}{R_2} \right) \quad (6.25)$$

Damit beträgt die gesamte Schwingungsdauer:

$$T = 2 \cdot t_1 = 2 \cdot R \cdot C \cdot \ln \frac{1 + \alpha}{1 - \alpha} = 2 \cdot R \cdot C \cdot \ln \left(1 + \frac{2 \cdot R_1}{R_2} \right) \quad (6.26)$$

Für $R_1 = R_2 = R$ vereinfacht sich die Gleichung zu:

$$T = 2 \cdot t_1 = 2 \cdot R \cdot C \cdot \ln 3 \approx 2,2 \cdot R \cdot C \quad (6.27)$$

7 Grundlagen digitaler Schaltungen

Lernziele

- Kennenlernen der grundlegenden Definitionen für digitale Schaltungen
- Struktur und Bezeichnungen von genormten Schaltzeichen nach DIN 40900

7.1 Inverter und Störabstände

In den analogen Schaltungen durfte die Ausgangsspannung die positive und negative Aussteuerungsgrenze nicht erreichen. In digitalen Schaltungen arbeitet man mit zwei Betriebszuständen. Man interessiert sich nur noch dafür, ob ein Spannungswert größer bzw. kleiner als ein vorgeschriebener Grenzwert ist. Die Grenzwerte werden als U_H für den „High“-Zustand und U_L für den „Low“-Zustand bezeichnet. Die exakten Zahlenwerte für diese Grenzwerte hängen von der eingesetzten Schaltungstechnik ab. Um eine eindeutige Zuordnung der Spannungen zu ermöglichen, sollen Spannungspegel zwischen U_L und U_H nicht vorkommen. Am Beispiel eines einfachen Pegelinverters in Emitterschaltung sollen diese Besonderheiten erklärt werden. Bild 7.1 zeigt die Schaltung und die Übertragungskennlinie. Die Schaltung soll folgende Funktion haben: für $U_e \leq U_{e,L}$ soll $U_a \geq U_H$ und für $U_e \geq U_{e,H}$ soll $U_a \leq U_L$ sein. Sperrt man den Transistor in Bild 7.1, wird die Ausgangsspannung im unbelasteten Fall $U_a = U_b$. Falls die niederohmigste Last $R_L = R_C$ ist, teilt sich die Spannung und es gilt $U_a = \frac{1}{2} \cdot U_b$. Sicherheitshalber definieren wir $U_{e,H} = 1,5 \text{ V}$. Das ist also die kleinste Eingangsspannung, bei der der Transistor voll eingeschaltet ist. Weiterhin soll die Eingangsspannung $U_e \leq U_{e,L}$ sein, damit $U_a \geq U_H$

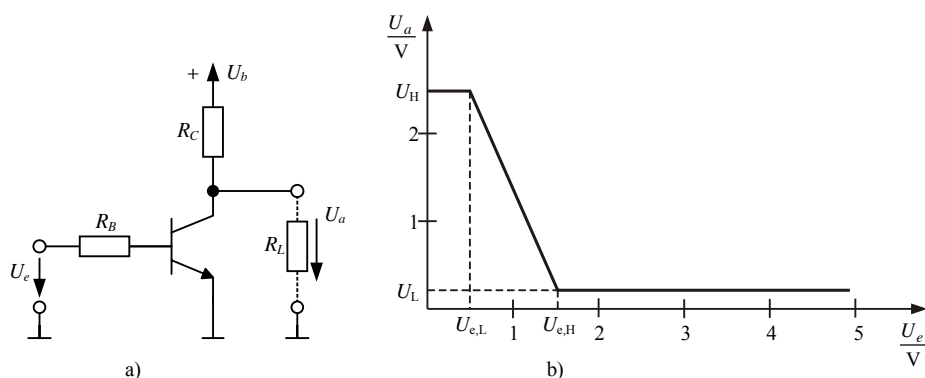


Bild 7.1: Emitterschaltung als Inverter (a) und Übertragungskennlinie für $R_V = R_C$ (b).

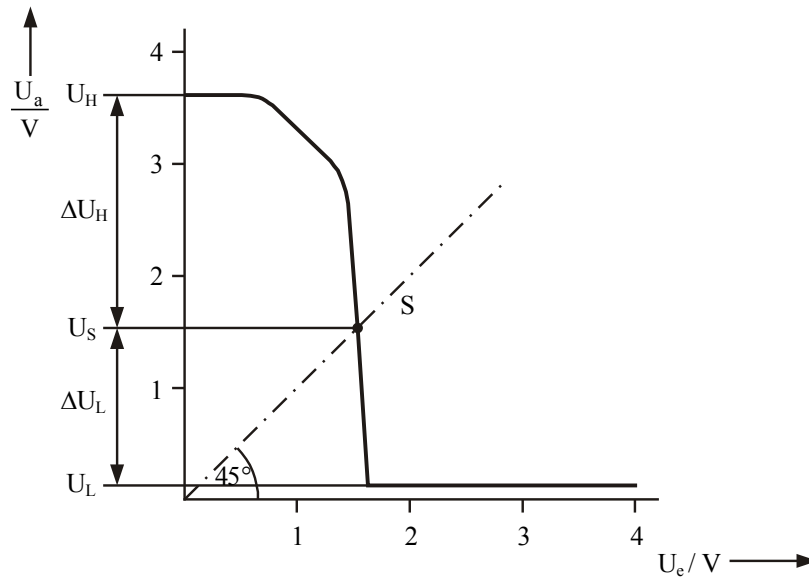


Bild 7.2: Übertragungskennlinie und Störabstände eines TTL-Inverters (7404).

wird. Als $U_{e,L}$ definieren wir die größte Eingangsspannung, bei der der Transistor gerade noch sperrt. Bei einem Si-Transistor sind das $U_{e,L} = 0,4 \text{ V}$. Der Einzelstörabstand gibt an, um wie viel die Eingangsspannung U_e eines Gatters in einer Kette von sonst ungestörten Gattern geändert werden darf, ohne dass die Information verfälscht wird. Dies ist der Fall, wenn der Schaltungspunkt S nicht überschritten wird. Am Beispiel der Übertragungskennlinie eines TTL-Inverters (Bild 7.2) können wir den Störabstand definieren. Dazu legen wir zunächst den Spannungshub $\Delta U = U_H - U_L$ fest. Für den H-Störabstand gilt:

$$\Delta U_H = U_H - U_S \quad (7.1)$$

Der L-Störabstand ist wie folgt definiert:

$$\Delta U_L = U_S - U_L \quad (7.2)$$

U_S ist dabei die Spannung im Schaltungspunkt der Übertragungskennlinie. Die Störabstände sind ein Maß für die Betriebssicherheit der Schaltung. Betrachten wir reale Störabstände bei einem Inverter des Typs SN7404. Bild 7.2 zeigt die Übertragungskennlinie des Gatters. Die relativen Störabstände können wie nachfolgend bestimmt werden:

$$Z_H = \frac{\Delta U_H}{\Delta U} \quad , \quad Z_L = \frac{\Delta U_L}{\Delta U} \quad (7.3)$$

Die in Bild 7.3 gezeigte Übertragungskennlinie des TTL-Inverters liefert mit $\Delta U = 3,65 \text{ V} - 0,1 \text{ V} = 3,55 \text{ V}$

die Einzelstörabstände:

$$\Delta U_L = 1,4 \text{ V und } \Delta U_H = 2,15 \text{ V}$$

die relativen Einzelstörabstände: $Z_L \approx 40 \%$ und $Z_H \approx 60 \%$

7.2 Anstiegs- und Gatterlaufzeiten

Realisierbare Impulse haben endlich steile Anstiegs- und Abfallsflanken. Da sie meistens auch nicht trapezförmig sind, definiert man die Anstiegs- und Abfallszeiten zwischen den 10 %- und 90 %-Punkten des Maximums, d.h. der Impulsamplitude nach Bild 7.3a. Die Symbole t_r und t_f stimmen mit den angelsächsischen Bezeichnungen für „rise“ und „fall“ time überein. Die Impulsdauer t_D wird zwischen den 50 %-Punkten gemessen. Gatterlaufzeiten werden meist durch andere physikalische Effekte bestimmt als Anstiegs- und Abfallszeiten. Beispielsweise kann eine lange verlustlose Leitung einen Rechteckimpuls praktisch unverzerrt, aber mit großer Verzögerung übertragen. In nichtlinearen Halbleiterschaltungen kann zudem der Übergang auf den niedrigeren Pegel verschieden stark gegenüber dem umgekehrten Übergang auf den höheren Pegel verzögert werden. In der Regel definiert man die Verzögerungen (propagation delay, High–Low, Low–High) t_{pdHL} und t_{pdLH} zwischen den 50 %-Punkten der Eingangs- und Ausgangsimpulse nach Bild 7.3. Die Gatterlaufzeit ist als Mittelwert aus den beiden Flankenverzögerungen definiert als

$$t_{pd} = \frac{t_{pdLH} + t_{pdHL}}{2} \quad (7.4)$$

Gelegentlich wird die Gatterlaufzeit auch als Mittelwert der Verzögerung zwischen den Schaltpunkten S nach Bild 7.3 angegeben, deren Spannungen U_S nicht den 50 % entsprechen.

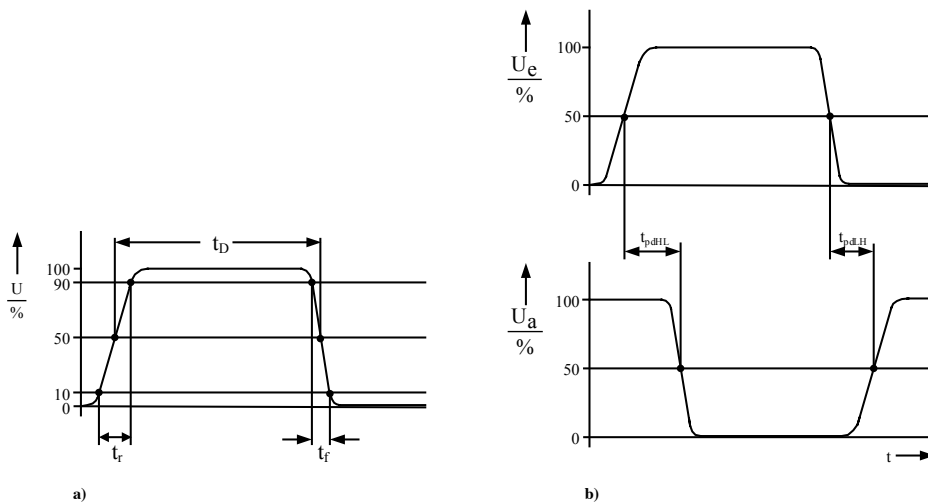


Bild 7.3: Charakteristische Zeiten von Impulsen: a) Impulsdauer t_D , Anstiegs- und Abfallzeit t_r und t_f und b) Verzögerung der Vorder- und Rückflanke t_{pdHL} und t_{pdLH} .

Die Unterschiede der Mittelwerte nach den beiden Definitionen sind jedoch meist vernachlässigbar. Die Gatterlaufzeit bipolarer Transistoren hängt von der Laufzeit der Minoritätsträger in der Basis τ_B und von der Zeitkonstante ab, die aus dem Lastwiderstand R_L und der Ersatzkapazität C_E berechnet werden kann

$$t_{pd} \approx k \cdot (\tau_B + R_L \cdot C_E) \quad (7.5)$$

wobei k eine Proportionalitätskonstante ist.

7.3 Die Verlustleistung

Integrierte Schaltungen haben eine endliche Verlustleistung P , die zur Erwärmung des Chips führt. Die Temperaturerhöhung des Chips beträgt bei einer umgesetzten Verlustleistung und einer Umgebungstemperatur T_A

$$\Delta T = T_{Chip} - T_A = P \cdot R_{th} \quad (7.6)$$

Dabei ist R_{th} der Wärmewiderstand zwischen Chip und Kühlkörper. Die maximale Verlustleistung, die von der Chipfläche ohne Zwangskühlung abgeführt werden kann, beträgt etwa 2 W/cm^2 . Der Momentanwert der Verlustleistung setzt sich aus einem statischen und dynamischen Anteil zusammen

$$P = P_{stat} + P_{dyn} \quad (7.7)$$

Die statische Verlustleistung ist durch die beiden quasistationären Betriebsströme I_1 und I_2 , die den beiden Spannungspegeln U_L und U_H entsprechen, definiert. Bei bekanntem Tastverhältnis $r = t_2/t_1$ und einer Taktfrequenz von $f_T = 1/t_p$ ($t_p = t_1 + t_2$) gilt

$$P_{stat} = [I_1 \cdot (1 - r) + I_2 \cdot r] \cdot (U_H - U_L) \quad (7.8)$$

Unter der Annahme von $I_1 = 0$ und $I_2 = (U_H - U_L)/R_L$ ergibt sich eine statische Verlustleistung nach

$$P_{stat} = \frac{r \cdot (U_H - U_L)^2}{R_L} \quad (7.9)$$

Die dynamische Verlustleistung ergibt sich aus den Umladungsprozessen der Lastkapazität. Die Lastkapazität setzt sich aus der Eingangskapazität der folgenden Gatter und der Kapazität der Verbindungsleitung zwischen den Gattern zusammen. Die dynamische Verlustleistung lässt sich näherungsweise wie folgt abschätzen. Die Energie, die in

einer Kapazität gespeichert ist, beträgt:

$$W = \frac{C \cdot U^2}{2} \quad (7.10)$$

Während einer Periode des Taktsignals wird die Kapazität zweimal umgeladen. Mit der Annahme, dass die Kapazität über eine ideale Stromquelle umgeladen wird, ergibt sich die mittlere Verlustleistung zu

$$P_{dyn} = C \cdot (U_H - U_L)^2 \cdot f_T \quad (7.11)$$

Die gesamte Verlustleistung beträgt somit

$$P = (r + R_L \cdot C \cdot f_T) \cdot \frac{(U_H - U_L)^2}{R_L} \quad (7.12)$$

Die Verlustleistung wächst mit zunehmender Taktfrequenz und Lastkapazität sowie abnehmendem Lastwiderstand. Ein oft benutztes Gütekriterium für das dynamische Verhalten logischer Schaltungen ist das Verzögerungszeit–Leistungs–Produkt (power–delay product). Darunter versteht man die Energie, die zum Schalten eines Gatters erforderlich ist:

$$W = P \cdot t_{pd} \quad (7.13)$$

Beide Faktoren dieses Produktes sollen möglichst klein sein.

7.4 Lastfaktoren

Jeder Gattereingang belastet den vorhergehenden Schaltkreisausgang weil ein bestimmter Strom zum Steuern des Gatters benötigt wird. Man definiert als Eingangslastfaktor 1 diejenige Belastung, die ein Eingang eines einfachen Grundgatters darstellt. So beträgt zum Beispiel bei der Transistor–Transistor–Logik (TTL) der Eingangsstrom bei L–Pegel $-1,6 \text{ mA}$ und bei H–Pegel $+40 \text{ }\mu\text{A}$. Der Eingangslastfaktor wird häufig als „Fan–in“ bezeichnet. Der Ausgangslastfaktor N_0 gibt an, mit welcher maximalen Anzahl von Gattereingängen der gleichen Schaltkreisfamilie mit einem Eingangslastfaktor 1 ein Ausgang belastet werden darf. Der Ausgangslastfaktor wird häufig als „Fan–out“ bezeichnet. Zum Beispiel haben TTL–Gatter in der Regel einen Ausgangslastfaktor von $N_0 = 10$. Leistungsgatter haben einen größeren Fan–out. Zusätzlich muss die kapazitive Last der Eingangskapazität des folgenden Gatters und die Leitungskapazität berücksichtigt werden. Bei CMOS–Schaltungen ist die maximal zulässige Ausgangsbelastung praktisch nur durch die kapazitive Last bestimmt, da der statische Eingangsstrom sehr gering ist.

7.5 Positive und negative Logik

Die Angaben für U_L und U_H sind keine logischen Zustände, sondern Spannungswerte. Diese beschreiben die Arbeitsweise der Schaltung. Welche logische Verknüpfung eine Schaltung erzeugt, kann erst gesagt werden, wenn die Spannungswerte den logischen Zuständen 0 und 1 zugeordnet worden sind. Die Zuordnung der Spannungspegel zu den logischen Zuständen kann folgendermaßen erfolgen:

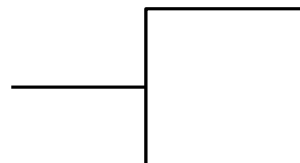
Positive Logik:	$U_L \Rightarrow 0$
	$U_H \Rightarrow 1$
Negative Logik:	$U_L \Rightarrow 1$
	$U_H \Rightarrow 0$

Heute wird in der Digitaltechnik vorwiegend mit positiver Logik gearbeitet.

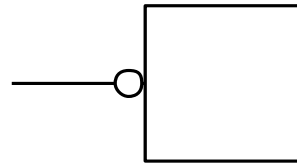
7.6 Genormte Schaltzeichen

Zur Darstellung logischer Funktionen digitaler Schaltungen werden vereinfachende Symbole verwendet, um komplexe Zusammenhänge durch Schaltpläne darstellen zu können. Schaltzeichen werden wie die Symbole einer Schrift im Laufe ihrer Geschichte auf den jeweiligen Stand der technischen Entwicklung durch die Festlegung von Normen angepasst. Die letzte Normung in der Bundesrepublik Deutschland fand 1984 statt und wurde unter der Bezeichnung DIN 40900 Teil 12 veröffentlicht. Sie enthält die internationale Norm IEC 617-12 von 1983. Auch wenn im englischsprachigen Raum nach wie vor andere Symbole verwendet werden, wollen wir uns hier mit den nach DIN genormten Schaltzeichen befassen. Die wichtigsten Grundelemente dieser Norm sind im Folgenden zusammengestellt und werden durch vereinfachte Formulierungen erläutert. Der genaue Wortlaut der Definitionen kann bei Bedarf in den DIN-Normen nachgelesen werden. Als Grundelement der logischen Symbole dient ein Rechteck, dessen Seitenkanten in den Abmessungen frei wählbar sind. Die logische Funktion und die Funktion der Ein- und Ausgänge werden durch zusätzliche Symbole definiert. Zuerst wollen wir uns mit der Darstellung logischer Eingänge befassen.

7.6.1 Darstellung von Eingängen

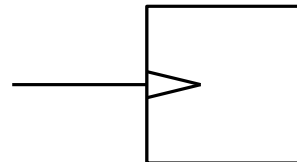


Nichtinvertierender Eingang (aktiv high)



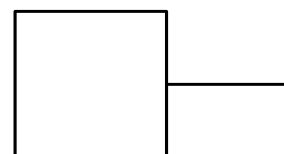
Invertierender Eingang (aktiv low)

Dynamischer Eingang. Hierbei findet eine interne Übernahme von logischer Information während der ansteigenden Flanke des Signals an diesem Eingang statt. Wird das Symbol mit dem vorigen (Kreis) kombiniert, geschieht dies bei der fallenden Flanke.



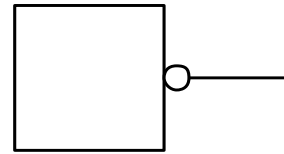
7.6.2 Darstellung von Ausgängen

Die Funktion von Ausgängen logischer Schaltungen ist recht umfangreich. Deshalb werden die Anschlüsse mit einer Reihe von verschiedenen Symbolen kombiniert. Wichtig ist hier, die Bedeutung der Symbole genau zu kennen, um klar zu unterscheiden, wie die logische Information am Ausgang zu werten ist, bzw. auf welche Art und Weise nachfolgende Bauelemente angeschlossen werden müssen. Im Gegensatz zu analogen Schaltungen, an denen wir entweder eine Spannung oder einem Strom als Ausgangsgröße erhalten, befassen wir uns bei logischen Schaltungen nur mit den logischen Pegeln H und L (oder 1 und 0). Die Ausgänge werden nun fast ausschließlich mit den Buchstaben Y (logische Gatter) und Q (Flipflops, Speicherschaltungen) gekennzeichnet. Die Ausgangssymbole besitzen folgende Bedeutung:

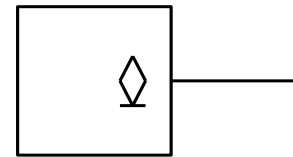


Nichtinvertierter Ausgang

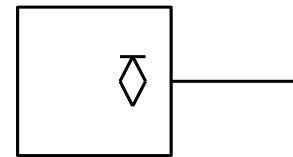
Invertierter Ausgang



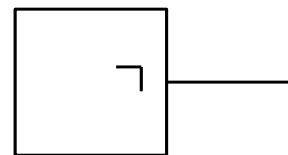
Offener Ausgang, aktiv Low, z.B. offener Kollektor eines npn Transistors, offener Drain bei n-Kanal Feldeffekttransistor. Im Low-Pegel ist der Ausgang niederohmig zum Masseanschluss.



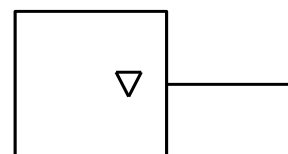
Offener Ausgang, aktiv High, z.B. offener Emitter eines npn Transistors, offener Source bei n-Kanal Feldeffekttransistor. Im High-Pegel ist der Ausgang niederohmig zum Betriebsspannungsanschluss.



Retardierter Ausgang (postponed output). Eine Zustandsänderung dieses Ausgangssignals wird so lange aufgeschoben, bis das Eingangssignal, welches die Zustandsänderung veranlasst hat, wieder in seinen ursprünglichen Ausgangszustand zurückgekehrt ist.



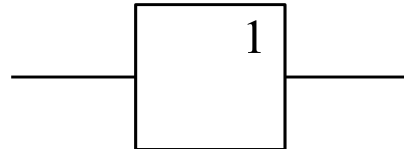
Tristate Ausgang. Zusätzlich zu den logischen Zuständen High und Low gibt es noch einen dritten, hochohmigen Zustand.



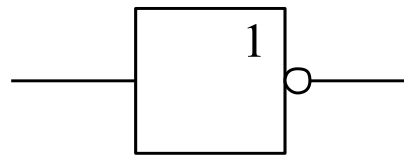
7.6.3 Logische Grundelemente

Die logische Funktion eines Gatter-Bausteins wird im Allgemeinen durch eine Bezeichnung in der rechten oberen Ecke des Symbols angegeben. Im Folgenden Abschnitt sollen zunächst nur die logischen Grundelemente betrachtet und erläutert werden.

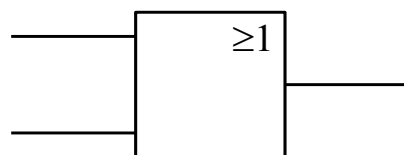
Treiber, Puffer. Am Ausgang liegt der gleiche logische Pegel, wie am Eingang



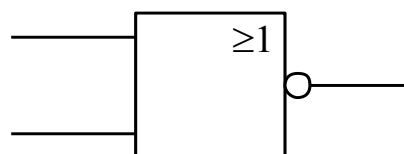
Inverter (NOT). Am Ausgang liegt das komplementäre Eingangssignal



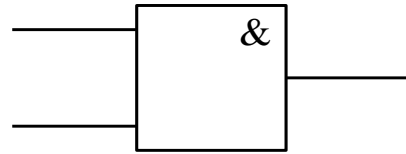
Oder-Gatter (OR)



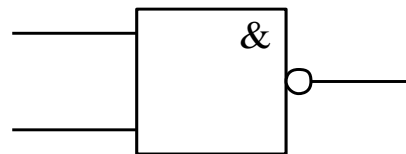
Nicht-Oder-Gatter (NOR)



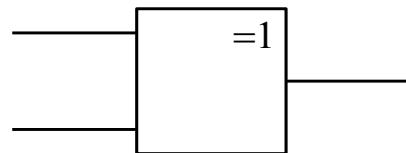
UND-Gatter (AND)



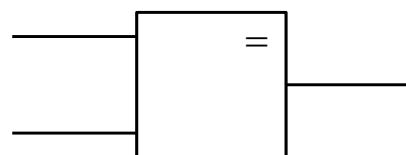
Nicht-UND-Gatter (NAND)



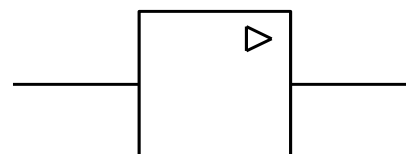
Antivalenz-Gatter (EXOR)

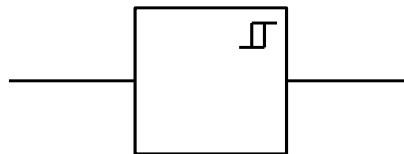


Äquivalenz-Gatter (EXNOR)

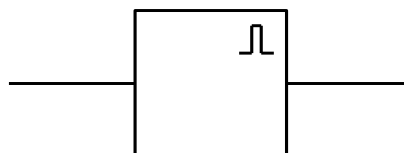


Leistungs-Element, Treiber (Buffer). Der Ausgangsstrom ist gegenüber normalen Logikbausteinen um ein Vielfaches höher.





Gatter mit Hysterese (Schmitt–Trigger)



Monostabiles Element (Monoflop)

Bevor wir uns mit Beispielen komplexerer Logikschaltungen befassen, müssen wir zuerst die Bedeutung von Bezeichnungen, die innerhalb der Symbole zu finden sind und Eingänge mit wichtiger zusätzlicher Information versehen, befassen. Erst mit dieser Bezeichnung ist eine exakte Zuordnung der Eingänge zueinander und zur logischen Funktion des Elements möglich. Für die Bezeichnung können Buchstaben und Ziffern verwendet werden. Grundsätzlich gilt folgende Regel:

Folgt eine Zahl einem Buchstaben, so handelt es sich um einen Eingang, der innerhalb der Logikschaltung mit mindestens einem weiteren Eingang verknüpft ist. Der oder die mit diesem Eingangssignal verknüpften anderen Eingänge führen dann die gleiche Ziffer vor dem Buchstaben.

Eine Bedingung ist immer dann erfüllt, wenn das Logiksignal im Innern des Elements einen High–Pegel bzw. eine 1 angenommen hat. Eingangsbezeichnungen können sein:

Cn	Der Clock- oder Takteingang gibt den mit der Zahl n gekennzeichneten Eingängen nur dann die definierte Wirkung, wenn Cn sich im internen 1-Zustand befindet.
Gn	Alle Eingänge, welche von diesem Eingang gesteuert werden, stehen in einer UND- Beziehung zu diesem Eingang.
nJ, nK, nD	Nimmt der J-Eingang den internen 1-Zustand an, wird im Element eine 1 gespeichert, nimmt der K-Eingang den internen 1-Zustand an, wird im Element eine 0 gespeichert. Sind beide Eingänge $J = K = 1$, so erfolgt bei jeder Periode des Cn-Signals ein Wechsel des Ausgangssignals in den komplementären Zustand. Der interne Logikzustand des D-Eingangs wird im Element gespeichert.
R	Nimmt der R-Eingang den internen 1-Zustand an, so wird im Element eine 0 gespeichert. (Reset-Eingang)
S	Nimmt der S-Eingang den internen 1-Zustand an, so wird im Element eine 1 gespeichert. (Set-Eingang) Die Eingänge R und S können auch als nR bzw. nS gekennzeichnet werden.
T	Nimmt der T-Eingang den internen 1-Zustand an, so wechselt der interne Zustand des Ausgangssignals in seinen komplementären Zustand. (Toggle-Eingang)

Manche Logikschaltungen besitzen neben einem Logikelement noch ein zusätzliches Steuerelement, in dem teilweise recht komplexe Funktionen untergebracht werden. Die in einem Steuerteil eines integrierten Bausteins verwendeten Buchstaben haben nach der Norm die Bedeutung einer Abhängigkeit der gesteuerten Eingänge von den Steuerungssignalen. Im Einzelnen sind dies:

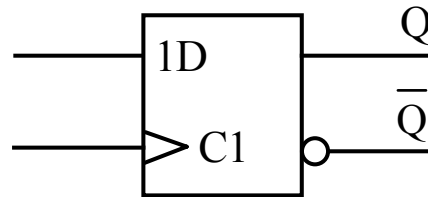
Gn	UND-Verknüpfung mit einem oder mehreren Eingängen
Vn	ODER-Verknüpfung mit einem oder mehreren Eingängen
Nn	EXOR-Verknüpfung (Antivalenz-Schaltung) mit einem oder mehreren Eingängen
Mn	Mode (Betriebsartenfestlegung)
ENn	Enable-Eingänge (Freigabesignal)
Ln	Load-Eingänge (Übernahmesignal für an bestimmten Eingängen anliegende Informationen)
Tn	Toggle-Eingänge
nS,nR	Set- und Reset Eingänge
Cn	Clock- oder Takt-Eingänge
An	Adressen-Eingänge

Anhand von Beispielen sollen die Kennzeichnungen der Eingänge nun näher erläutert werden.

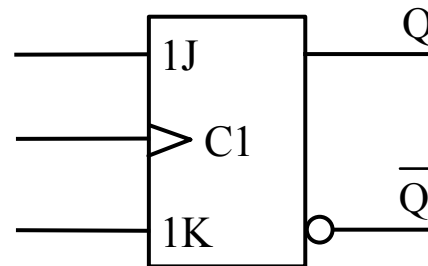
7.6.4 Beispiele

1. Bistabile Elemente:

Durch die Kennzeichnung mit zwei Ausgängen, die eine nichtinvertierte und eine invertierte logische Information anzeigen, ist üblicherweise zu erkennen, dass ein bistabiles Logikelement enthalten ist. Hier handelt es sich um ein D-Flipflop, dessen logischer Pegel am Eingang 1D mit der ansteigenden Flanke des Signals am Eingang C1 übernommen, und sofort an den Ausgang weitergeleitet wird.

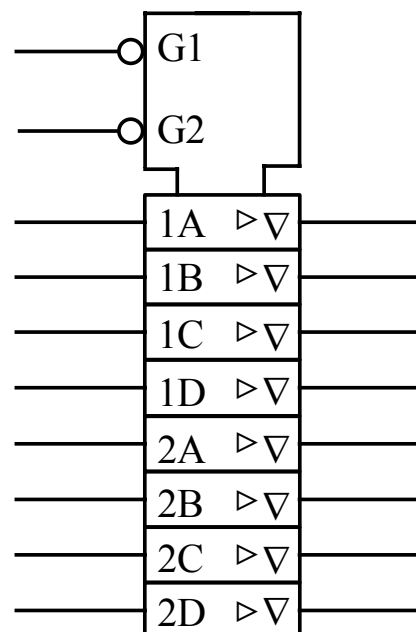


Hier handelt es sich um ein JK-Flip-Flop, dessen logische Pegel am Eingang 1J bzw. 1K mit der ansteigenden Flanke des Signals am Eingang C1 übernommen, und sofort an den Ausgang weitergeleitet werden. Man spricht auch davon, dass mit der ansteigenden Flanke das Bauelement getriggert wird.



2. Komplexe Bausteine:

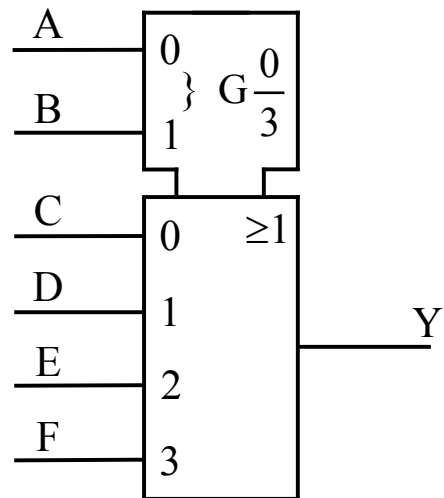
Hier wird ein sogenannter Bustreiber-Baustein mit Tristate-Ausgang dargestellt. Er enthält ein Steuerelement und acht Logikelemente. Der Steuereingang G1 ist mit den Eingängen 1A - 1D und der Eingang G2 mit 2A - 2D durch eine UND-Funktion verknüpft. Liegt am Eingang G1 ein LOW-Pegel an, so werden die entsprechenden logischen Pegel an den Eingängen 1A - 1D an die zugehörigen Ausgänge durchgeschaltet. Entsprechendes gilt für G2 und die Eingänge 2A bis 2D. Liegt an den Steuereingängen ein HIGH-Pegel an, sind die Ausgänge hochohmig.



4 zu 1 Multiplexer: Das Zeichen $G_{\frac{0}{3}}$ sagt aus, dass die Eingänge 0 und 1 als binäre Adresse der Eingänge C bis F aufzufassen sind und mit einem von diesen eine UND-Funktion bilden. Beispielsweise ist für $A = 1$, $B = 0$ der Eingang D mit dem Ausgang Y 'verbunden'.

Wahrheitstabelle:

B	A	Y
0	0	C
0	1	D
1	0	E
1	1	F



8 Schaltkreisfamilien

Lernziele

- Kennenlernen des Aufbaus der unterschiedlichen Schaltkreisfamilien für integrierte Digitalschaltungen
- Schaltkreisfamilien mit bipolaren Transistoren
- Schaltkreisfamilien in CMOS-Technik
- Verstehen von elektronischen Schaltern mit Feldeffekt-Transistoren

8.1 Bipolare integrierte Schaltungen

Logische Gatter bilden die elementaren Grundbausteine digitaler Schaltungen und Systeme. Die Bezeichnung Gatter (Tor) weist darauf hin, dass diese „Tore“ durch den Einfluss der an den Eingängen liegenden Spannungen geöffnet oder geschlossen werden können. Auf diese Art und Weise können Informationen weitergeleitet oder ihre Weiterleitung verhindert werden. Die logische Funktion der Gatter, d.h. die logische Verknüpfung von Aus- und Eingängen, wird mit Hilfe der Schaltalgebra beschrieben. Die meisten industriell hergestellten Schaltungen enthalten als Grundelemente NAND- bzw. NOR-Gatter. Diese Gatter sind ökonomisch günstig herstellbar und universell einsetzbar. Aus der Vorlesung „Grundlagen der Digitaltechnik“ sind die grundlegenden Funktionen der Schaltalgebra und ihre Realisierung mit Hilfe von Schaltwerken bekannt. Relaisschaltungen sind als einfache Schaltwerke einsetzbar. Sie werden heute noch in der Starkstromtechnik eingesetzt. Komplexere logische Verknüpfungen erfordern jedoch universell einsetzbare und schnelle Gatter. Diese Gatter können mit Dioden, bipolaren Transistoren oder MOS-Transistoren aufgebaut werden. Gatter, die nach bestimmten Prinzipien aufgebaut werden, nennt man eine Schaltkreisfamilie. Gatter einer Schaltkreisfamilie lassen sich problemlos zu komplexen Verknüpfungsschaltungen zusammenschalten, da sie einheitliche Spannungspegel für HIGH- und LOW-Pegel haben.

8.1.1 DTL-Schaltungen

Wird die logische Verknüpfung der Eingangssignale mit Dioden vorgenommen, erhält man die Dioden-Transistor-Logik (DTL). Bild 8.1 zeigt ein NAND-Gatter mit zwei Eingängen. Liegt an den Eingängen ein HIGH-Pegel an, sind die Eingangsströme aufgrund der in Sperrrichtung betriebenen Eingangsdioden gering. Liegt an den Eingängen ein LOW-Pegel (L-Pegel) an, wird der Strom in den Ausgangstransistor des vorherigen Gatters durch den Widerstand R_1 bestimmt. Bei Gattern mit mehreren Eingängen

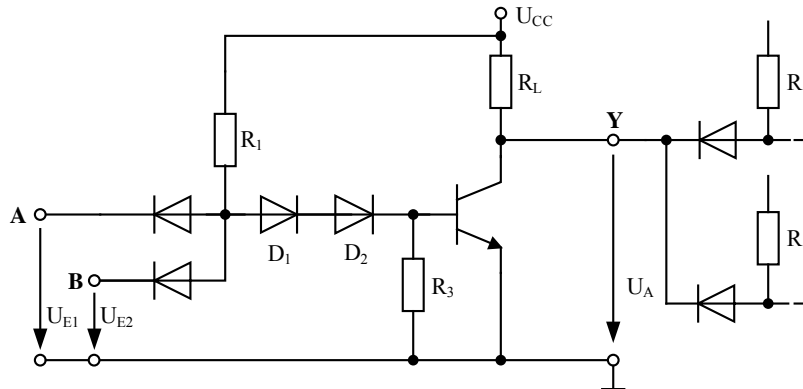


Bild 8.1: Schaltbild eines DTL-NAND-Gatters mit zwei Eingängen.

addieren sich die Ströme, so dass am Widerstand R_1 eine recht hohe Verlustleistung entstehen kann. Der Widerstand R_3 dient zum schnellen Ableiten der Basisladung des Ausgangstransistors beim Ausschalten. Ersetzt man die Diode D_1 der DTL-Schaltung nach Bild 8.1 durch einen Transistor und einen weiteren Widerstand R_2 (siehe Bild 8.2a), so wird der Eingangsstrom I_E für einen L-Pegel am Eingang um den Faktor $R_1/(R_1+R_2)$ verkleinert und damit zu:

$$I_E = -\frac{U_{CC} - U_D - U_{CE}}{R_1 + R_2} \quad (8.1)$$

U_D ist die Durchlassspannung der Eingangsdiode und U_{CE} die Restspannung des vollständig eingeschalteten Ausgangstransistors der vorhergehenden Stufe. Bild 8.2a zeigt diese verbesserte Version der DTL-Schaltung für eine Versorgungsspannung $U_{CC} = 5 \text{ V}$. Durch die Reduzierung des Eingangsstroms wird auch die Verlustleistung im Gatter deutlich reduziert. Wird anstelle der Siliziumdiode D_2 zwischen Eingangs- und Ausgangstransistor eine Zener-Diode mit $U_Z = 6,8 \text{ V}$ eingesetzt und die DTL-Schaltung bei einer Versorgungsspannung von $U_{CC} = 15 \text{ V}$ betrieben (Bild 8.2b), erhält man einen wesentlich größeren Störabstand und eine nahezu symmetrische Übergangskennlinie um den Schaltpunkt wie Bild 8.2c zeigt.

8.1.2 TTL-Schaltungen

Im Gegensatz zur DTL-Schaltung wird die logische AND-Verknüpfung der Eingangssignale eines Gatters bei der TTL-Technologie (Transistor-Transistor-Logik) mit einem Multi-Emitter-Eingangstransistor durchgeführt. Multi-Emittertransistoren können besonders platzsparend integriert werden, da die Emitterdiffusionsbereiche nebeneinander ohne zusätzliche pn-Isolationsringe im gleichen Basisdiffusionsbereich angebracht werden. Eine besonders einfache Auslegung einer TTL NAND-Schaltung mit einem Multi-Emitter-Transistor am Eingang und einem Transistor mit „offenem“ Kollektor am Ausgang zeigt Bild 8.3. Um die logischen Spannungspegel am Ausgang aufzubauen, muss ein externer Lastwiderstand R_L zwischen der Versorgungsspannung und dem Kollektor des

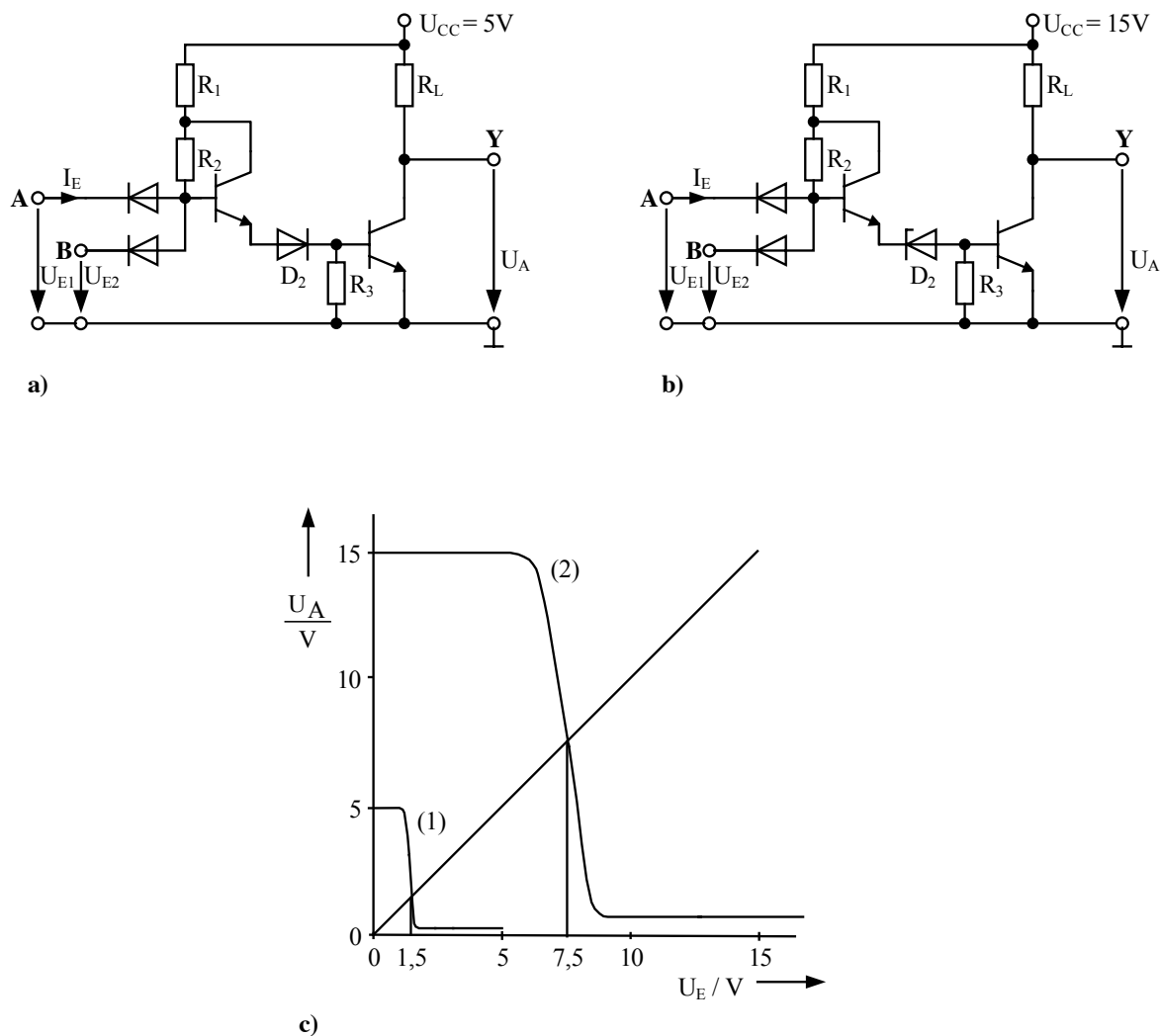


Bild 8.2: DTL-NAND-Gatter mit zwei Eingängen und erheblich kleineren Eingangsströmen als nach Bild 8.1, b) DTL-NAND-Schaltung mit großem Störabstand, c) Übertragungskennlinien: (1) normale pn Diode D_2 nach Bild 8.2a, (2) Zener-Diode nach Bild 8.2b.

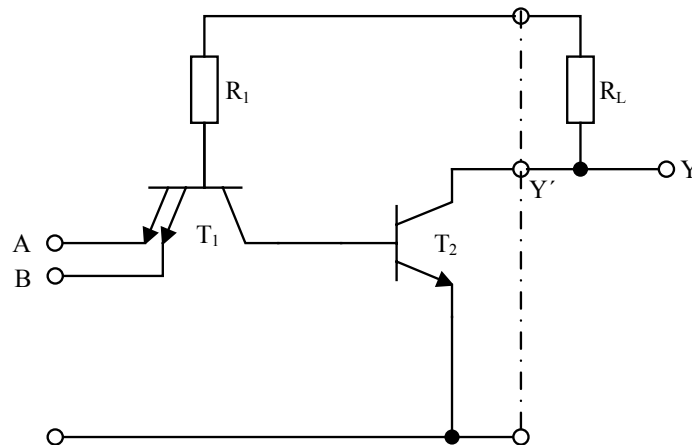


Bild 8.3: TTL-NAND-Grundsaltung mit npn Transistoren und offenem Kollektor.

Ausgangstransistors vorhanden sein. Wird an einen der beiden Eingänge A oder B ein L-Pegel und an den anderen ein H-Pegel angelegt, ist der Transistor T_1 eingeschaltet. An der Basis des Transistors T_2 liegt dann ebenfalls ein L-Pegel an, so dass T_2 sperrt. Der Kollektor des Transistors T_2 , der als Ausgang der Schaltung dient und über einen externen Widerstand R_L mit U_{CC} verbunden ist, liegt damit auf dem H-Pegel. Der Basisstrom von T_2 fließt durch den auf L-Pegel liegenden Gatter-Eingang aus dem Schaltkreis heraus. Die Summe aller Emitterströme ist gleich dem Basisstrom von T_2 , der sehr klein ist. Wird nun an beide Eingänge A und B ein H-Pegel angelegt, so wird die Spannung am Emitter von Transistor T_1 höher als die Spannung am Kollektor ($U_{BE(T_2)}$) und der Transistor wird invertiert betrieben (Inversbetrieb). Der Basis-Kollektor pn-Übergang ist in Durchlassrichtung gepolt. Über die Basis-Kollektordiode von Transistor T_1 fließt in Durchlassrichtung ein Basisstrom in den Transistor T_2 und schaltet diesen ein. Am Ausgang Y stellt sich ein L-Pegel mit $U_{CE,sat}$ (Kollektor-Emitter-Sättigungsspannung) ein. Die Entladung einer Kapazität am Ausgang Y kann sehr schnell erfolgen, da der Widerstand R_{CE} zwischen Kollektor und Emitter des Transistors T_2 im eingeschalteten Zustand sehr klein wird. Damit der L-Pegel einen möglichst niedrigen Wert annimmt, muss zugleich das Verhältnis R_L/R_{CE} möglichst groß sein. Andererseits führt ein zu großer Widerstand R_L auf zu große Zeitkonstanten $R_L \cdot C$ beim Aufladen der Lastkapazität C am Ausgang Y der Schaltung. Auch würde die Spannung U_H des H-Pegels stark von der Zahl der angeschlossenen Gatter abhängig werden:

$$U_H = U_{CC} - n \cdot I_E \cdot R_L \quad (8.2)$$

Wobei n die Anzahl der Eingänge ist, und I_E der Eingangsstrom eines Eingangs. Durch die Erweiterung mit einer Emitterfolgerstufe werden die Eingangsströme noch kleiner und die Änderung der Eingangsspannung verstärkt, d.h. die Übertragungskennlinie fällt steiler ab. Die Gegentaktendstufe mit niedrigem Ausgangswiderstand ermöglicht eine wesentlich schnellere Umladung der Lastkapazität, als dies vorher mit dem Lastwi-

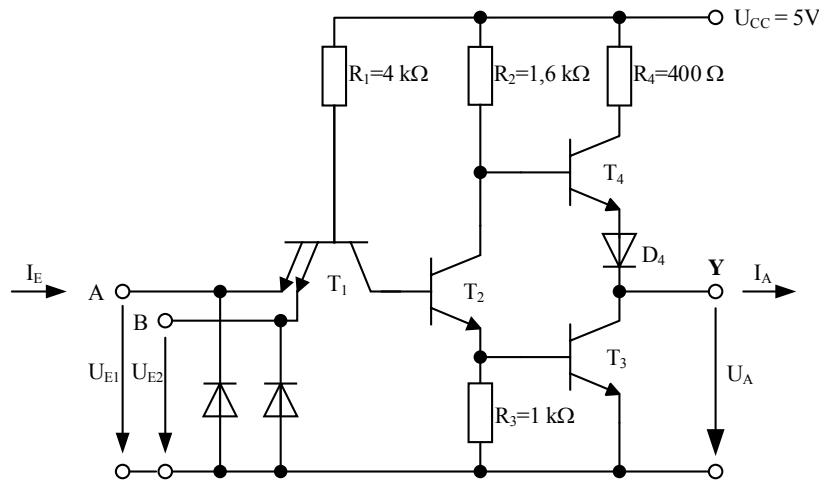


Bild 8.4: Schaltbild einer Standard TTL-NAND-Schaltung.

derstand R_L möglich war. Diese Standard-TTL-Schaltung eines NAND-Gatters mit Eingangsschutzdioden und der Gegentakt-Endstufe der sogenannten „Totem-Pole Ausgangsstufe“ ist in Bild 8.4 dargestellt. Im Gegensatz zu der in Bild 8.3 dargestellten Schaltung benötigt man keinen externen Lastwiderstand, um beide logische Pegel definiert einzustellen. Liegt an einem der beiden Eingängen A oder B eine Spannung $U_E = 0$ V an, leitet der Transistor T_1 . Die Transistoren T_2 und T_3 sind gesperrt und Transistor T_4 ist leitend, so dass der Ausgang der Schaltung sich im High-Zustand befindet. Die Endstufe kann einen relativ großen Ausgangsstrom liefern, der dann die Lastkapazität schnell auflädt. Liegt an beiden Eingängen eine hohe Eingangsspannung $U_{E1} = U_{E2} = 5$ V an, sind T_2 und T_3 eingeschaltet (in der Sättigung), und T_4 sperrt. Die Lastkapazität kann damit über T_3 schnell entladen werden.

Gatter mit Tristate Ausgang

Bei Anwendungen, in denen TTL-Schaltkreise an Bussysteme angeschlossen werden, ist es zweckmäßig, Schaltkreise mit Tristate Ausgang einzusetzen. Bei diesen Schaltkreisen kann der Ausgang durch ein Steuersignal in den hochohmigen Zustand geschaltet werden. Der Ausgang dieser Schaltkreise kann somit drei Zustände annehmen: Low- und High-Zustand mit einem niederohmigen Ausgangswiderstand und einen hochohmigen Zustand. Den hochohmigen Zustand realisiert man durch Sperren der beiden Ausgangstransistoren. Bild 8.5 zeigt die Schaltung eines TTL-Gatters mit Tristate-Ausgang. Beim Anlegen eines H-Pegels an den Steuereingang S werden die Transistoren T_7 und T_8 leitend. Das bewirkt, dass der dritte interne Eingang des Multi-Emitter-Transistors T_1 auf Null gezogen wird, wodurch T_2 und T_5 gesperrt werden. Zusätzlich erniedrigt D_1 die Basisspannung von T_3 auf etwa 0,7 V, so dass auch T_3 und T_4 gesperrt sind. Der Ausgang des Gatters ist somit hochohmig. Das Steuersignal S wird auch „output enable“-Signal genannt. Beim Anlegen eines L-Pegels an S wird T_6 leitend und damit T_7 und T_8 gesperrt. Die Schaltung arbeitet als normales TTL-Gatter

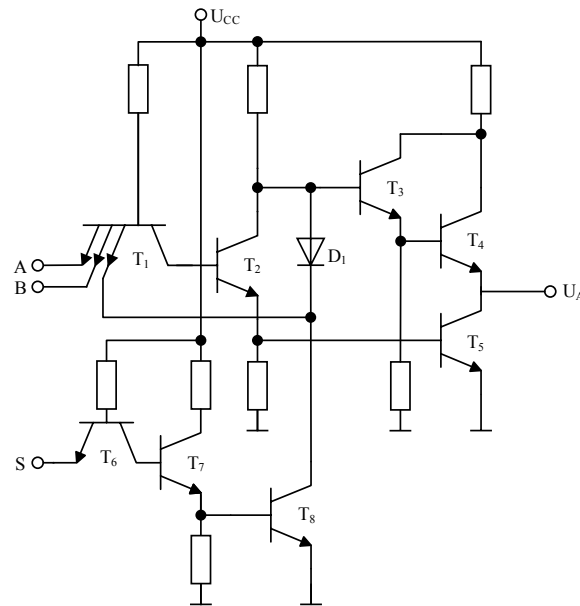


Bild 8.5: Schaltbild einer TTL-Schaltung mit Tristate-Ausgang.

mit zwei Eingängen. Der Ausgang ist aktiv, d.h. er kann L- oder H-Pegel annehmen. Standard-TTL-Schaltungen haben eine große Gatterlaufzeit, da die Transistoren in die Sättigung umschalten.

8.1.3 Schottky- und Low-Power Schottky-TTL-Schaltungen

Zur Verwirklichung kurzer Gatterlaufzeiten muss neben den Abfall- und Anstiegszeiten vor allen Dingen die Speicherzeit klein gehalten werden. Die „Speicherzeit“ beim bipolaren Transistor ist ein Effekt, der durch die Forderung nach einem schnellen Einschalten des Transistors und einer kleinen Kollektor-Emitterrestspannung hervorgerufen wird. Beides erfordert einen relativ großen Basisstrom. Wird die Eingangsspannung U_E sprunghaft von 0 V auf U_{CC} erhöht, wird auch sprunghaft ein Basisstrom i_B fließen, der im Einschaltmoment aufgrund der Sperrschichtkapazität der Basis-Emitterdiode überschwingt (Bild 8.6b). Der Transistor wird leitend, d.h. die Kollektor-Emitterspannung U_{CE} sinkt vom H-Pegel auf den L-Pegel ab. Die Basis-Emitterspannung beträgt $U_{BE} = 0,7$ V. Erreicht die Kollektor-Emitterspannung Werte unter 0,7 V, wird nun auch die Basis-Kollektordiode leitend. Der Kollektor kann jetzt nicht mehr die vom Emitter injizierten Ladungsträger aus der Basis aufnehmen; vielmehr werden auch noch vom Kollektor Ladungsträger in die Basis emittiert. Damit befinden sich viel mehr freie Ladungsträger in der Basis, als notwendig.

Beim Abschalten des Transistors muss jetzt zuerst die in der Basis vorhandene überschüssige Ladung abgebaut werden, bevor die Kollektor-Emitterspannung wieder 0,7 V erreicht, und der Kollektor wieder als Kollektor arbeiten und damit der Transistor mit dem Abschalten beginnen kann. Die dazu notwendige Zeit nennt man Speicherzeit. Der

Transistor hat also eine wesentlich größere Verzögerungszeit beim Abschalten als beim Einschalten ($t_{pdLH} \gg t_{pdHL}$). Der Eintritt in die Sättigung kann durch einen großen Basisvorwiderstand R_B nur bis zu einem gewissen Grade vermieden werden, da im praktischen Betrieb zur Gewährleistung des unteren Spannungspegels, trotz Exemplarstreuungen, Schwankungen der Temperatur, der Versorgungsspannung und verschiedener Belastungen am Ausgang, die Kollektor-Emitter-Spannung klein werden muss. Eine weitere Möglichkeit besteht darin, parallel zum Basiswiderstand einen Kondensator zu legen. Durch den sehr kleinen Blindwiderstand des Kondensators beim Schaltvorgang werden zwar beim Ein- und Ausschalten die Abfall- bzw. Anstiegszeit der Ausgangsspannung reduziert, aber mit dem Nachteil, dass große Widerstände und Kapazitäten integrierter Schaltungen eine große Chipoberfläche beanspruchen und damit für höhere Integrationsdichten nicht geeignet sind.

Käufliche integrierte Digitalschaltungen besitzen deshalb kein lineares RC-Glied am Eingang, sondern eine nichtlineare Schottky-Diode zur Begrenzung der Basis-Kollektor-Spannung. Die Schaltung eines Schottky-Dioden-Inverters und die zugehörigen Strom- und Spannungsverläufe sind in Bild 8.6 zu sehen.

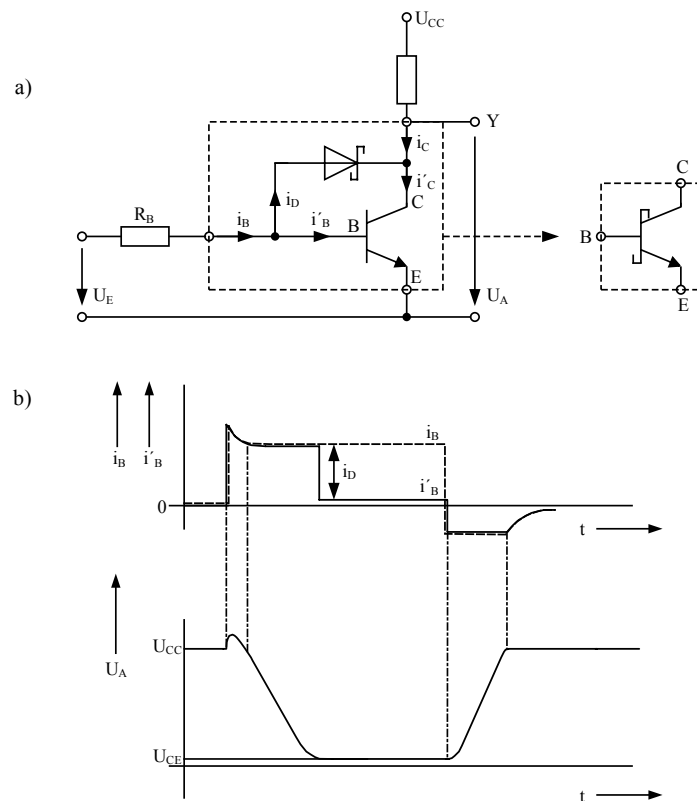


Bild 8.6: (a) Emitterschaltung mit diskreter Schottky-Diode zwischen Basis und Kollektor zur Vermeidung der Sättigung, Symbol eines Transistors mit integrierter Schottky-Diode, (b) Zeitabhängigkeit der Ströme und Spannungen beim Ein- und Ausschalten.

Nach Anlegen einer sprungförmigen positiven Spannung U_E übernimmt zunächst der Transistor den vollen Basisstrom $i_B = i'_B$, da die Ausgangsspannung noch größer als die Eingangsspannung, d.h. $U_A > U_E$ ist, und daher die Schottky-Diode gesperrt bleibt. Wenn die Ausgangsspannung unter $U_{CE} = 0,7 \text{ V}$ absinkt, werden gleichzeitig die Schottky- und die Basis-Kollektordiode in Durchlassrichtung betrieben. Die Schottky-Diode übernimmt aber den überschüssigen Basisstrom, da sie mit $U_{KSD} \approx 0,3 \text{ V}$ eine um etwa $0,4 \text{ V}$ kleinere Knickspannung als die pn-Diode des Transistors zwischen Basis und Kollektor besitzt, so dass letztere weit unterhalb ihrer Knickspannung betrieben wird. Der Basisstrom i_B des Schottky-Dioden-Transistors kann nach Bild 8.6b wesentlich größer gewählt werden als für den quasistationären Kollektorstrom notwendig wäre ($i_B \gg i'_B$). Der L-Pegel der Ausgangsspannung erhöht sich aber um die Knickspannung der Schottky-Diode. Damit wird die Sättigung der Basis vermieden. Beim Ausschalten kann deshalb der Basisstrom des Schottky-Transistors praktisch ohne Verzögerung der Spannung folgen, da sich keine überschüssigen Ladungsträger in der Basis befinden, also auch kein Speichereffekt existiert. Die Ausgangskennlinienfelder eines Transistors mit und ohne Schottky-Diode sind in Bild 8.7a und b gezeigt. Eine $I - U$ Kennlinie einer einzelnen Schottky-Diode ist in Bild 8.7c zu sehen. Die beiden Transistorkennlinien unterscheiden sich im Wesentlichen nur durch eine um etwa $U_{KSD} = 0,3 \text{ V}$ größere Kollektor-Emitter-Restspannung des Schottky-Dioden-Transistors, da eine Aussteuerung bis tief in die Sättigung, $U_{BC} > 0,3 \text{ V}$, verhindert wird. Eine Schottky-Diode kann mit weniger als 10 % zusätzlicher Chipfläche integriert werden ohne dass Zuleitungskapazitäten auftreten. Die Schottky-Transistor-Logik hat gegenüber der gesättigten Transistorlogik die folgenden Vorteile:

- 1.) wesentlich kürzere Schaltzeiten durch Vermeidung der Sättigung,
- 2.) kleinere Verlustleistung durch kleinere Spannungshübe, kleinere Umschaltenergien,
- 3.) höhere Stromverstärkung, da eine Gold-Dotierung nicht mehr gebraucht wird, um durch verstärkte Rekombination die Sättigungszeit zu reduzieren,
- 4.) flexiblerer Schaltungsentwurf, da pn- und Schottky-Dioden sowie npn-Transistoren zur Verfügung stehen.

Diese Vorteile erreicht man auf Kosten:

- 1.) eines ausgefeilteren und aufwendigeren Entwurfs insbesondere für besonders schnelle Schaltungen,
- 2.) geringerer technologischer Herstellungstoleranzen,
- 3.) eines kleineren Störabstandes infolge einer größeren Kollektor-Emitter-Restspannung.

Die Schottky-Dioden-Technik hat sich heute für die mittleren und höheren Geschwindigkeitsbereiche weitgehend durchgesetzt. Bild 8.8 zeigt eine einfache Auslegung eines Schottky-TTL-Gatters, die praktisch identisch mit der TTL-Schaltung in Bild 8.4 ist.

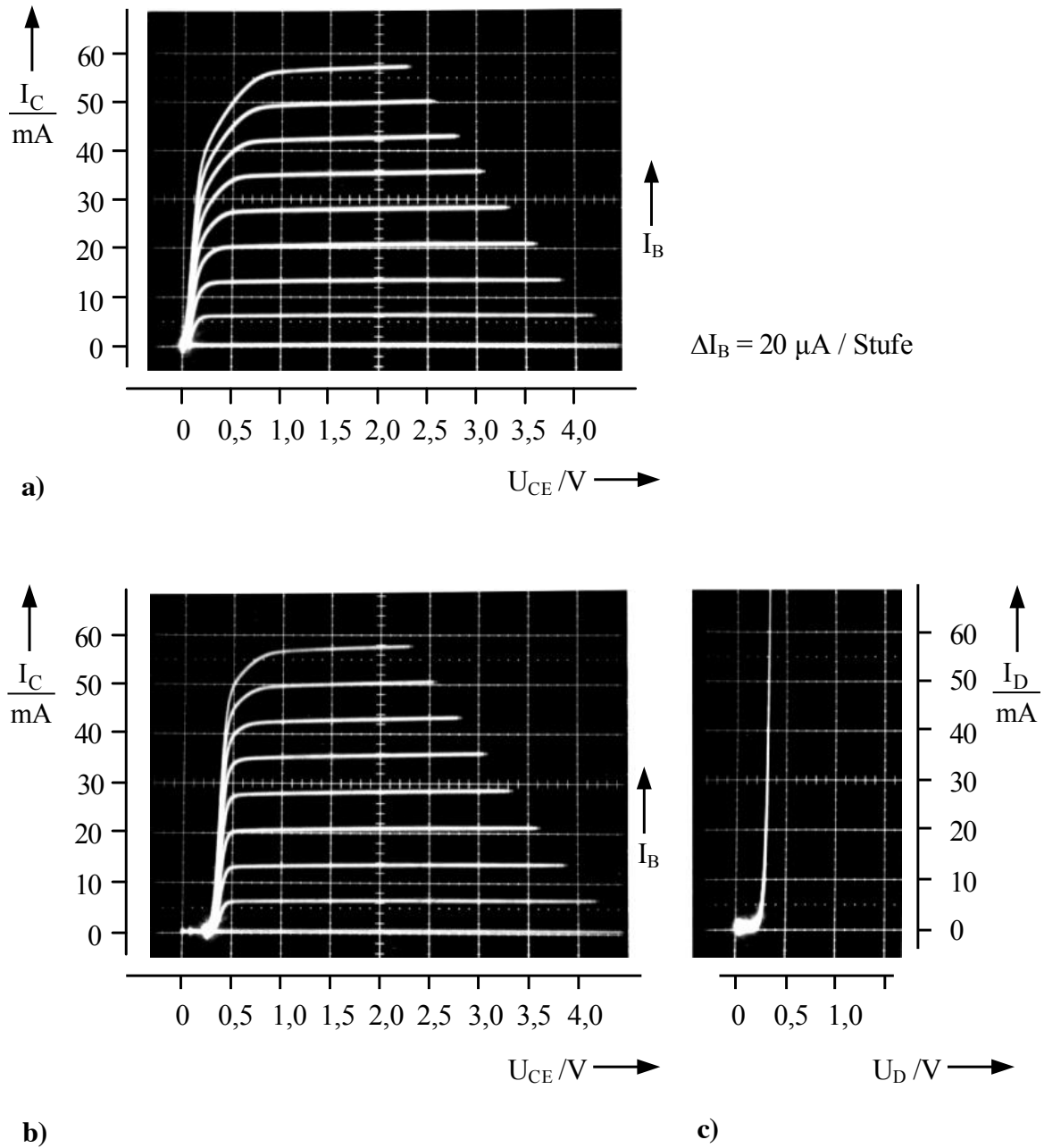


Bild 8.7: Kennlinienfelder des Transistors und der Schottky-Diode der Inverterschaltung in Bild 8.6: a) Transistor ohne Schottky-Diode, b) Transistor mit Schottky-Diode, c) Kennlinie der Schottky-Diode zwischen Basis und Kollektor

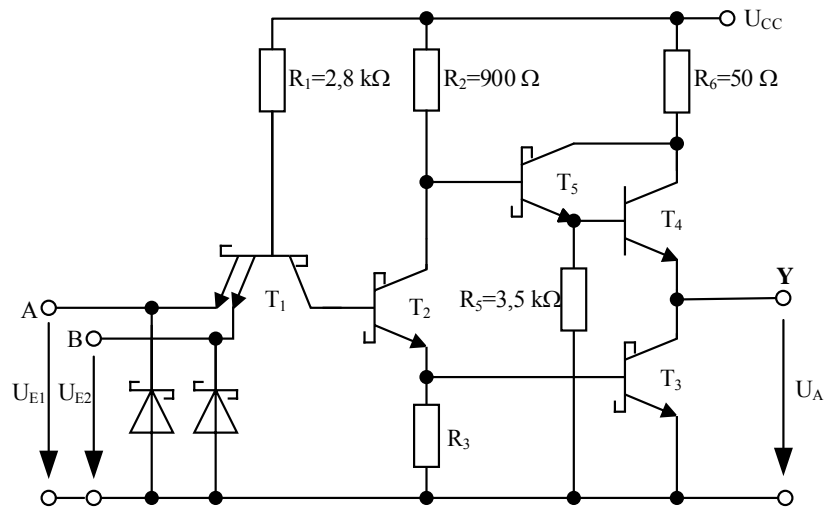


Bild 8.8: Schottky-TTL-NAND-Gatter mit Schottky-Schutzdioden am Eingang.

Logik-Familie	t_{pd} / ns	P_{stat} / mW	R_L / Ω
Gesättigte Technik: SN 74 00	10	10	400
Ungesättigte Technik: SN 74 LS 00	10	2	120
SN 74 S 00	3	19	50
SN 74 ALS 00	4	1	50
SN 74 AS 00	1,7	8	20

Tabelle 8.1: Übersicht über bipolare Logikfamilien. Es bedeuten: LS: Low-Power Schottky; S: Schottky; ALS: Advanced Low-Power Schottky; AS: Advanced Schottky.

Low-Power-Schottky-TTL

Durch den Einsatz der Schottky-Diode wird die Sättigung der Transistoren vermieden. Das führt zu einer relativ hohen Verlustleistung der Schaltkreise. Die Eingangsschaltung der Low-Power Schottky-Schaltungen (LS-TTL) weichen deutlich von der TTL-Grundschialtung ab. Sie basieren auf der Diodenverknüpfung, wie sie bereits bei der DTL-Familie eingesetzt wurde. Bild 8.9 zeigt die Schaltung eines LS-TTL-Gatters. Die beiden Eingänge sind mit Schottky-Dioden zur logischen Verknüpfung und einer Schutzschaltung vor negativen Überspannungen ausgelegt. In Tabelle 8.1 sind die typischen Daten eines NAND-Gatters mit zwei Eingängen der Baureihe 74.. von Texas Instruments für die wichtigsten TTL-Schaltkreistechnologien zusammengefasst. Die angegebenen Gatterlaufzeiten gelten für eine Lastkapazität von $C_L = 15$ pF

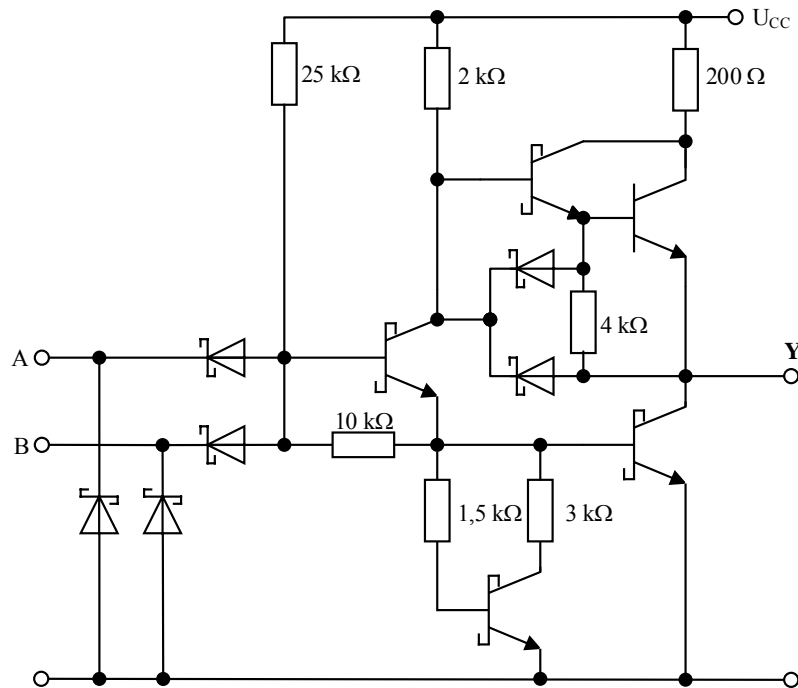


Bild 8.9: Low-Power Schottky-TTL-NAND-Gatter mit Schottky-Schutzdioden am Eingang.

8.1.4 Übersicht über die Schaltungsfamilien

Eine übersichtliche Gegenüberstellung der Schaltungsfamilien mit ihren wichtigsten Kenndaten ist in Tabelle 8.2 gegeben. Eine Zusammenstellung der bipolaren Logikschaltkreise ist in Tabelle 8.3 dargestellt.

	Eingangspegel / V				Ausgangspegel / V			
	Low		High		Low		High	
	min.	max.	min.	max.	min.	max.	min.	max.
DTL	–	1,4	3,6	–	–	0,5	4,0	–
TTL	–	0,8	2,0	–	–	0,4	2,4	–
LS-TTL	–	0,8	2,0	–	–	0,4	2,4	–
NMOS	–0,5	0,65	2,2	5	0	0,45	2,4	5
CMOS	0	2	3	5	0	0,01	4,99	5

Tabelle 8.2: Übersicht über die Schaltpegel der verschiedenen Schaltungsfamilien.

Schaltkreisfamilie	Typische Grundschaltung eines Gatters (vereinfacht)	Vor- und Nachteile	Bemerkungen zur Funktion	Bemerkungen zum Einsatz in der Informationstechnologie (IT)
DTL – (Dioden-Transistor-Logik)		+ schneller und störlicher als RTL + gute Entkopplung der Eingänge - relativ großer Flächenbedarf in IS (Isolationsinseln ...)	• logische Verknüpfung durch D_1 und D_2 • mittels D_3 wird der Störabstand verbessert	• früher häufige Anwendung beim Aufbau aus Einzelelementen
TTL – (Transistor-Transistor-Logik / SN 7400)		+ nur eine Betriebsspannung + sehr gut für IS geeignet + niederohmiger Ausgang ($< 100 \Omega$) + Serie von internationalen Firmen hergestellt (austauschbar) - Einsatz in der AT z.T. problematisch (Störfähigkeit) - Relationen bei längeren Verbindungsleitungen ($> 25..50 \text{ cm}$) - kurzer Stromstoß (10 bis 15 mA, 6..10 ns) auf der Betriebsspannungsführung beim Umschalten des Gatters	• logische Verknüpfung durch Multitransistor (Standard-TTL) • es gibt 7 Bauformen der TTL-Familie: Low-Power-Schottky-TTL, Standard-TTL, High-Speed-TTL, ALS-TTL, AS-TTL	• in Folge niedriger Signalpegel und hoher Schaltgeschwindigkeit wesentlich störfähiger als frühere Logikschaltungen aus Einzelbauelementen. Deshalb z.T. Zusatzaufwand beim Einsatz in der IT erforderlich

Tabelle 8.3: Übersicht der bipolaren Logikschaltkreise

8.2 MOS–Schaltkreise

Die MOS-Technologie hat gegenüber der bipolaren Technologie nachfolgende Vorteile, die ihren Einsatz für die Herstellung von integrierten Schaltkreisen bestimmen:

- Geringer Bedarf an Chipfläche
- Einfache Herstellungstechnologie
- Kleine Verlustleistung
- Leistungslos steuerbar

Gegenwärtig haben die n-Kanal-MOS-, CMOS- und BiCMOS-Technologie eine praktische Bedeutung für ICs.

8.2.1 NMOS–Schaltungen

Schaltet man einen n-Kanal MOS-Transistors vom Anreicherungstyp über einen ohm'schen Widerstand R_{DD} an die Betriebsspannung U_{DD} , so entsteht eine wichtige Grundsaltung der Digitaltechnik (siehe Bild 8.9a). Diese Schaltung entspricht der Source-Grundsaltung, die in Kapitel 3.4 besprochen wurde. Bei der Eingangsspannung $U_E = 0$ V am Gate ist der Transistor gesperrt und die Ausgangsspannung hat den maximalen Wert U_{DD} . Umgekehrt ist die Ausgangsspannung sehr klein, wenn die Eingangsspannung einen großen Wert hat, z.B. $U_{DD} = 10$ V. Diese Grundsaltung hat die Funktion eines Inverters. Man spricht von einem logischen Inverter, wenn die Spannung am Eingang nur zwei Pegel annehmen kann, z.B. 0, 3 V und 5 V oder 10V, die häufig mit L (LOW) und H (HIGH) gekennzeichnet werden. Solche Inverter können auf einfache Weise logisch hintereinander geschaltet werden, d.h. die Ausgangsspannung des einen Inverters kann direkt ohne zusätzliche Schaltmaßnahmen mit dem Eingang des folgenden Inverters gekoppelt werden.

In Bild 8.10 sind drei Möglichkeiten für den Aufbau eines Inverters in der NMOS Technik und deren Übertragungskennlinien gezeigt. Bild 8.10a zeigt einen Inverter mit einem ohm'schen Widerstand als Last und dessen Übertragungskennlinie. Beim Durchlaufen der Eingangsspannung U_E von Null bis U_{DD} gelangt der Schalttransistor vom Sperrbereich über den Sättigungsbereich bis in den linearen Bereich. Da die Herstellung von Widerständen sehr teuer ist und viel Chipfläche benötigt, ist es einfacher einen Lasttransistor aus der gleichen technologischen Bauart nach Bild 8.10b einzusetzen bei dem Gate und Drain miteinander verbunden sind ($U_{GS} = U_{DS}$). Der Lasttransistor arbeitet hier als nichtlinearer Widerstand. Für Eingangsspannungen, die den Transistor sperren, bleibt ein Spannungsabfall über dem Lasttransistor. Somit ist die Ausgangsspannung im H-Pegel um etwa U_{th} geringer als bei einem ohm'schen Lastwiderstand. Dieser Nachteil wird vermieden, wenn als Lasttransistor anstelle eines Anreicherungstyps ein Verarmungstyp mit einer negativen Schwellspannung eingesetzt wird, bei dem Source und Gate verbunden bleiben, also $U_{GS} = 0$ V ist. Dieser Fall ist in Bild 8.10c

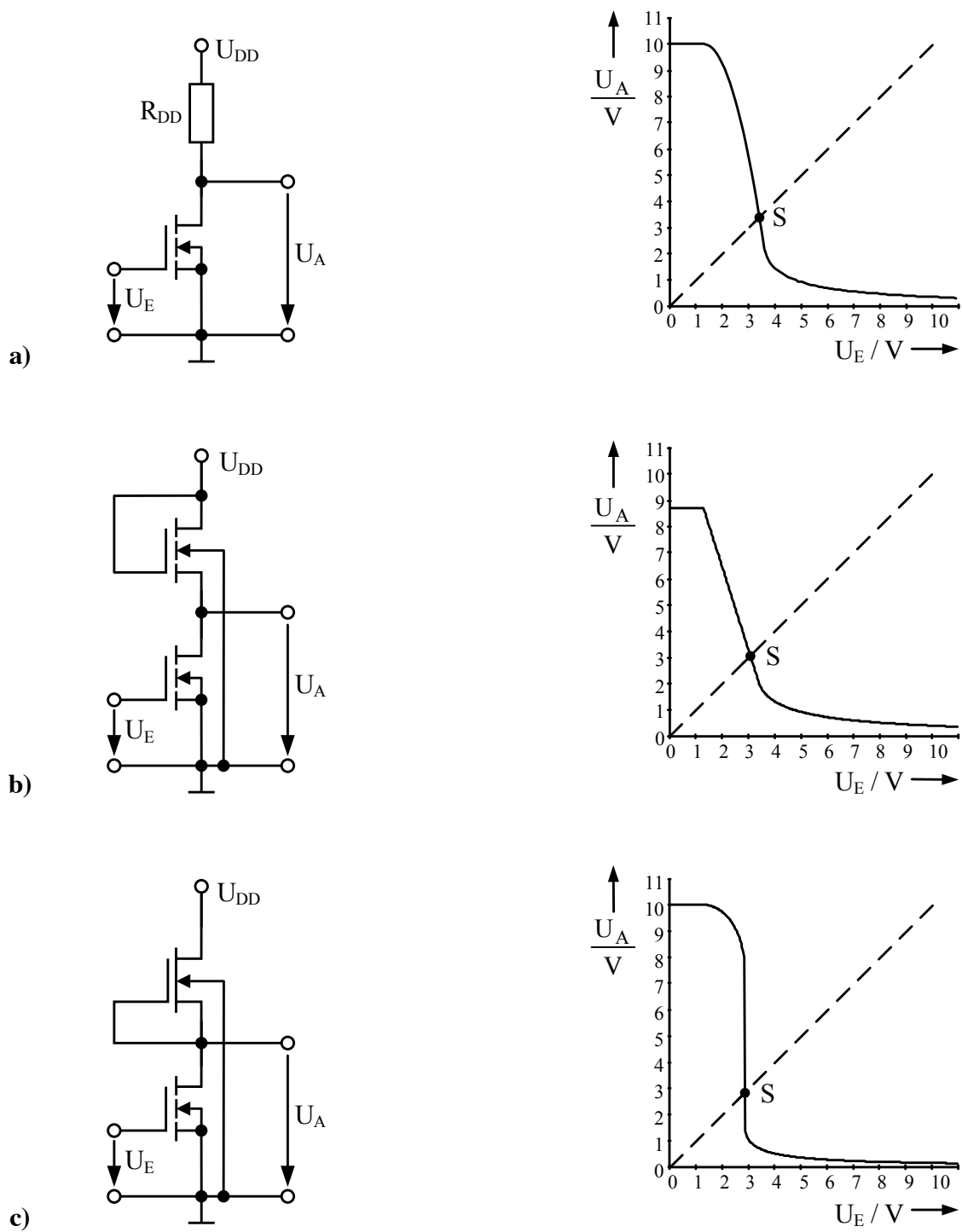


Bild 8.10: Schaltbilder und Übertragungskennlinien eines Inverters a) mit Lastwiderstand R , b) mit Lasttransistor vom Anreicherungstyp (AT) und c) mit Lasttransistor vom Verarmungstyp (VT).

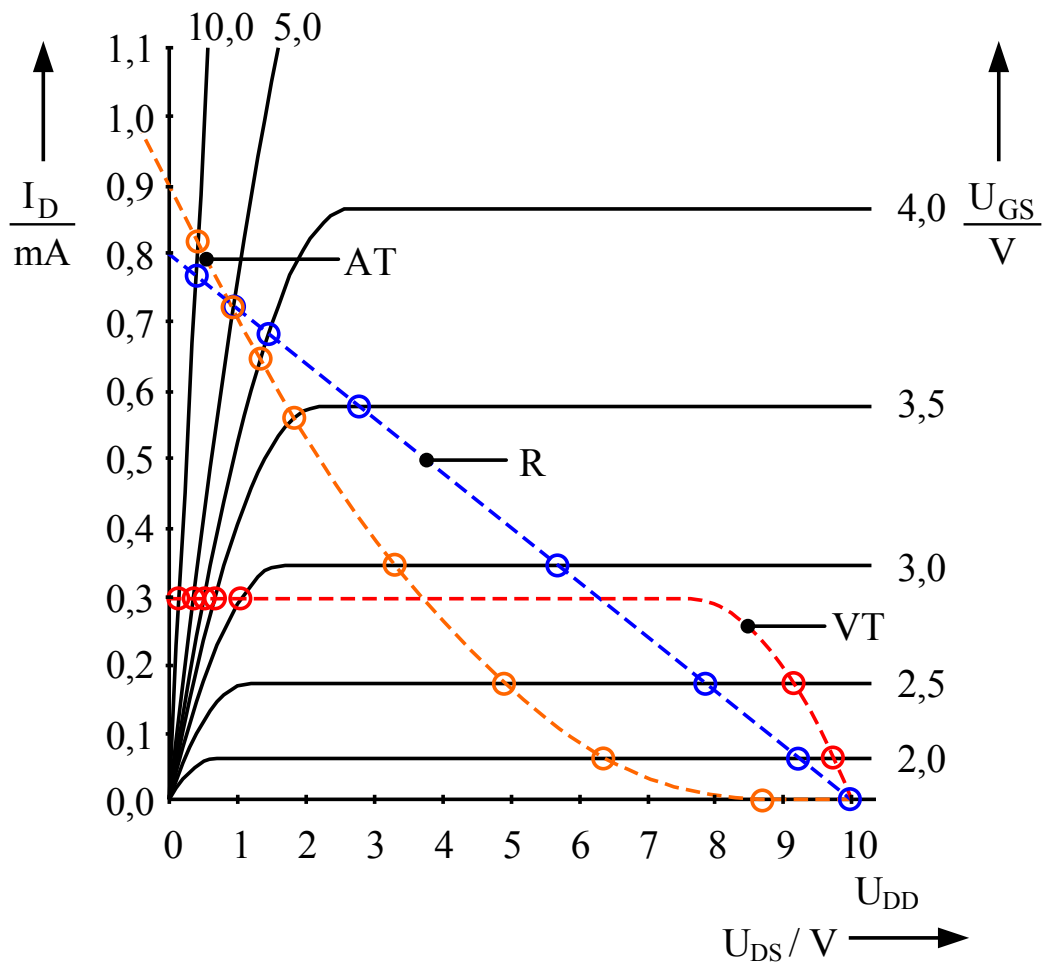


Bild 8.11: Kennlinienfeld des Schalttransistors mit den Lastfunktionen der drei Inverter nach Bild 8.10a-c.

dargestellt. Die volle Betriebsspannung steht wie bei einem linearen Lastwiderstand als Ausgangsspannung zur Verfügung. Der Lasttransistor vom Verarmungstyp wirkt in einem weiten Spannungsbereich wie eine Stromquelle. Der maximale Strom und die Verlustleistung ($I^2 \cdot R_L$) sind erheblich kleiner durch den Einsatz eines Lasttransistors vom Verarmungstyp anstelle eines Transistors vom Anreicherungstyp. Die Übergangskennlinie ist in Bild 8.10c eingezeichnet. Die Lastkennlinien für den ohm'schen Widerstand und die beiden Transistoren sind in Bild 8.11 in das Kennlinienfeld des Schalttransistors eingezeichnet. Die Lastkennlinien für den ohm'schen Widerstand und die beiden Transistoren sind in Bild 8.11 in das Kennlinienfeld des Schalttransistors eingezeichnet.

Verknüpfungsglieder mit NMOS-Schaltungen

Bild 8.12 zeigt eine NOR- und NAND-Schaltung in NMOS-Technik mit einem ohm'schen Arbeitswiderstand. Ein NOR-Gatter entsteht dadurch, dass man die Ausgänge

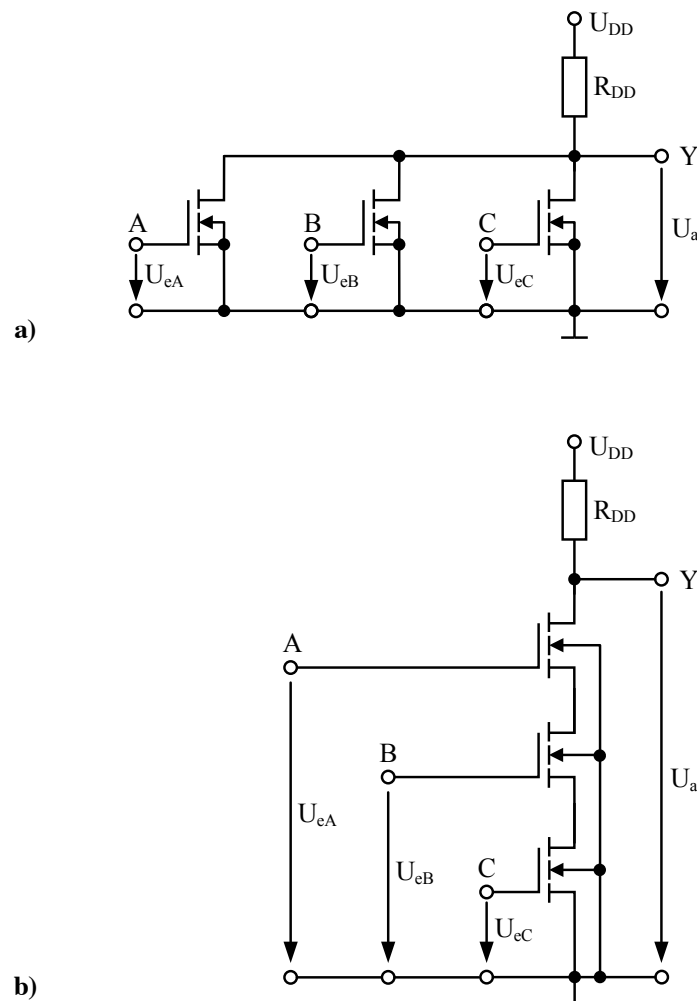


Bild 8.12: a) NOR- und b) NAND-Schaltung mit n-Kanal MOSFETs vom Anreicherungstyp.

mehrerer Inverter parallel schaltet und nur ein Lastwiderstand für alle Inverter eingesetzt wird. Wird einer der drei Eingänge mit dem H-Pegel angesteuert, so schaltet der Ausgang auf den L-Pegel. Zur Realisierung einer NAND-Schaltung werden die Eingangstransistoren in Reihe geschaltet. Liegt an allen drei Eingängen H-Pegel an, so schaltet der Ausgang auf LOW. Bild 8.13 zeigt die Schaltungen von NAND- bzw. NOR-Gattern mit Lasttransistor.

8.2.2 CMOS-Schaltungen

Aufbau und Funktionsweise eines CMOS-Grundgatters wurden im Kapitel 3.5 besprochen. Die Bilder 8.13 und 8.14 zeigen entsprechend die Schaltung und Übertragungsfunktion einer CMOS-Grundsaltung. Die Übertragungsfunktion ist bei der Hälfte der Versorgungsspannung sehr steil und verläuft symmetrisch zum Umschaltpunkt. Das ein-

fache Grundgatter bildet einen Inverter, der nochmals in Bild 8.14a dargestellt ist. Die Übertragungskennlinien des CMOS-Inverters für verschiedene Betriebsspannungen sind in Bild 8.14b gezeigt.

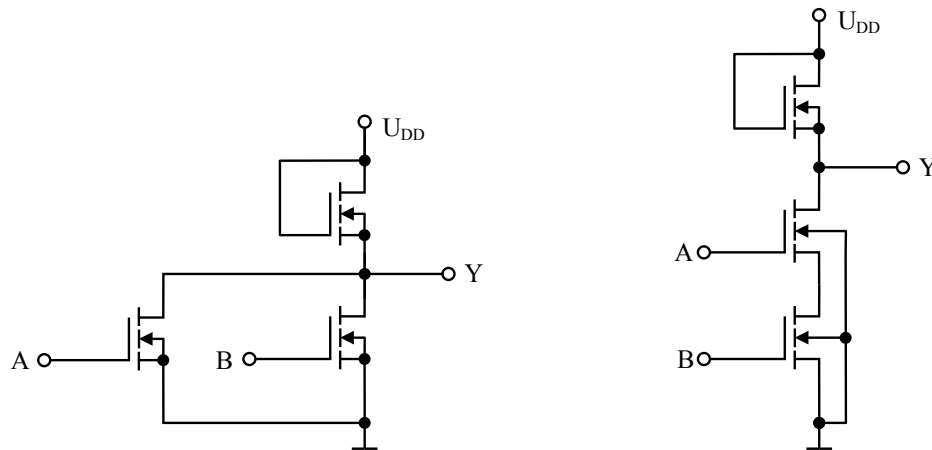


Bild 8.13: NOR- und NAND-Schaltung mit n-Kanal MOSFETs vom Anreicherungstyp mit Lasttransistor vom Anreicherungstyp.

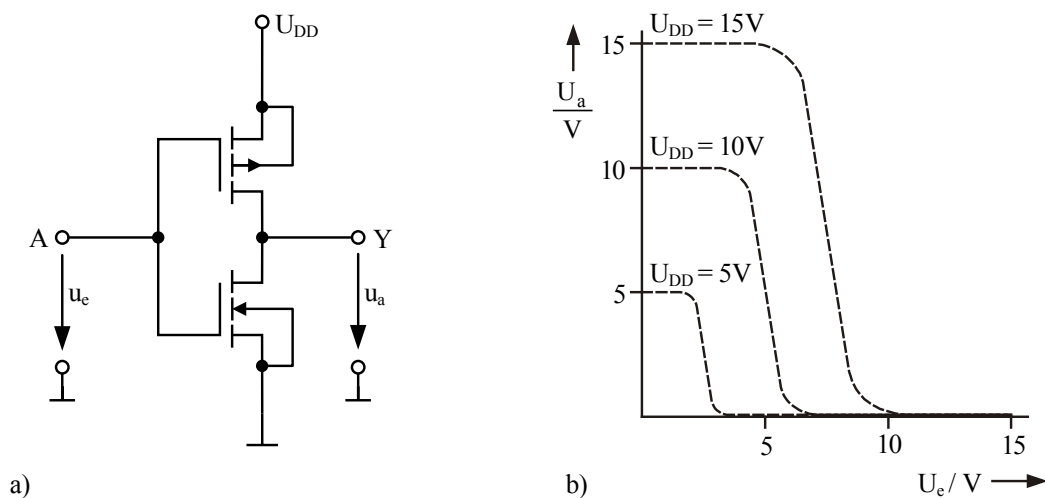


Bild 8.14: a) Schaltung eines CMOS-Inverters und b) Übertragungskennlinien für verschiedene Betriebsspannungen.

Eine NAND-Schaltung in CMOS-Technologie, die 3 Eingangssignale verknüpft, ist in Bild 8.15 zu sehen. Nur wenn alle Eingänge A bis C eine positive Spannung führen, die dem H-Pegel entspricht, sind die drei in Reihe geschalteten n-Kanal Transistoren leitend, so dass der Ausgang Y den L-Pegel annimmt. Da in diesem Fall die drei parallel geschalteten p-Kanal Lasttransistoren sperren, fließt bis auf Leckströme kein Strom

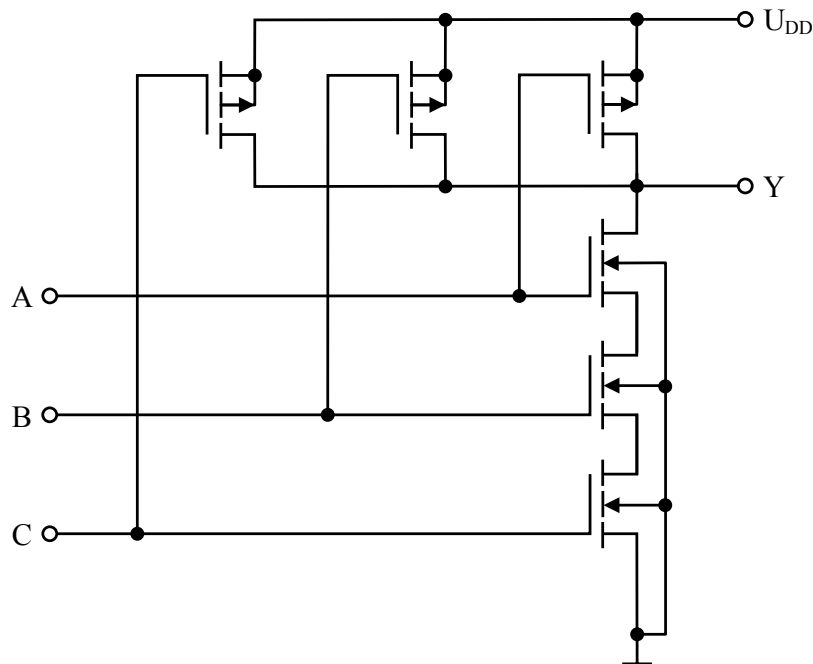


Bild 8.15: Schaltung eines NAND-Gatters in CMOS-Technik.

über die NAND-Schaltung, wenn der L- oder H-Zustand am Ausgang erreicht worden ist. Große Ströme und Verlustleistungen treten kurzzeitig während des Umschaltens auf, wenn beispielsweise A und B auf der Versorgungsspannung U_{DD} liegen und C den Bereich um $U_{DD}/2$ durchläuft. Die mittlere Verlustleistung der CMOS-NAND-Schaltung bleibt also wie bei einem CMOS-Inverter außerordentlich klein. Wird der Ausgang Y der Schaltung in Bild 8.15 durch große Eingangskapazitäten C der folgenden Stufen belastet, müssen die Transistoren entsprechend groß ausgelegt werden um die notwendigen Ströme zum Auf- und Entladen der Lastkapazitäten aufzubringen. Die Aufladezeitkonstante τ ist dem Durchlasswiderstand eines Lasttransistors der NAND-Schaltung proportional. Daher kann die Halbleiteroberfläche der NAND-Schaltung miniaturisiert werden, wenn die Lastkapazität klein bleibt. Dem Ausgang von CMOS-Grundgattern sind nur kleine Ströme entnehmbar, die gerade zur Ansteuerung von wenigen CMOS Eingängen oder einem Eingang von Low-Power (LS)-TTL-Schaltungen ausreichen.

Um den Ausgangsstrom zu erhöhen, werden Pufferstufen eingesetzt. Bild 8.16 zeigt ein NAND-Gatter mit zwei nachgeschalteten Invertern als Ausgangspuffer. Ein NAND-Gatter in Bild 8.16 erzeugt als Funktion der Eingangssignale am Ausgang entweder eine niederohmige Verbindung mit der Masse oder der Versorgungsspannungsleitung.

Sie sind also nicht geeignet für den direkten Anschluss an die Leitung eines Bussystems mit vielen Sendern, die zu verschiedenen Zeiten und an verschiedenen Orten einspeisen sollen. Daher werden Schaltungen benötigt, deren Ausgang während der Sendepause

unabhängig von den logischen Eingängen hochohmig ist. Man spricht von Tristate-Schaltungen, wenn neben den beiden logischen auch ein hochohmiger Zustand eingestellt werden kann. Ein Tristate-Inverter ist in Bild 8.17 skizziert. Liegt seine Steuerspannung U_{St} an Masse, so wird der Ausgang Y durch je einen gesperrten Transistor von Masse und Versorgungsspannung getrennt. Für $U_{St} = H$ ist die normale Inverterfunktion gewährleistet.

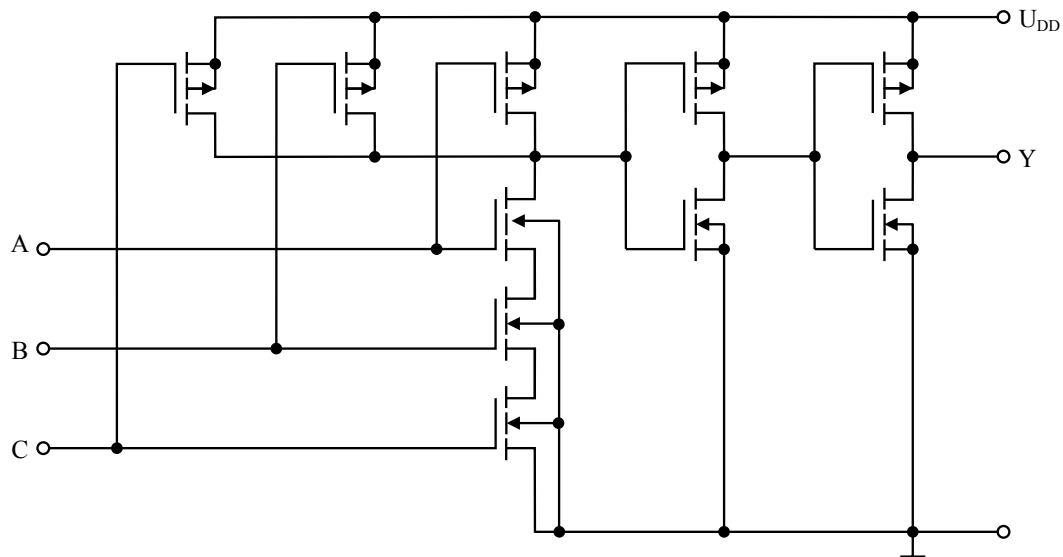
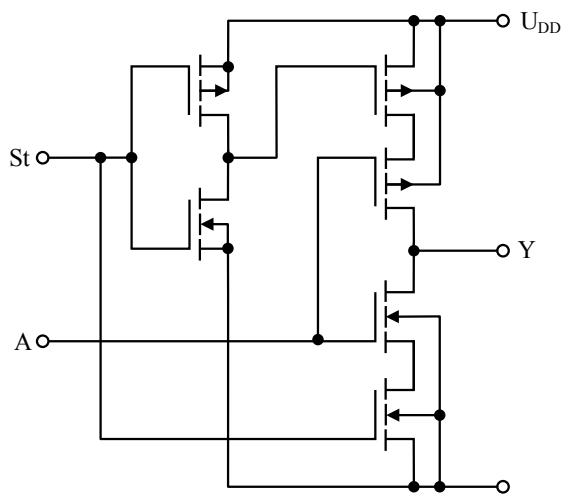


Bild 8.16: Schaltung eines NAND-Gatters mit Treiberstufen in CMOS-Technik.



a)

St	A	Y
0	X	HiZ
1	0	1
1	1	0

b)

Bild 8.17: Tristate CMOS-Inverter. a) Schaltung, b) Wahrheitstabelle

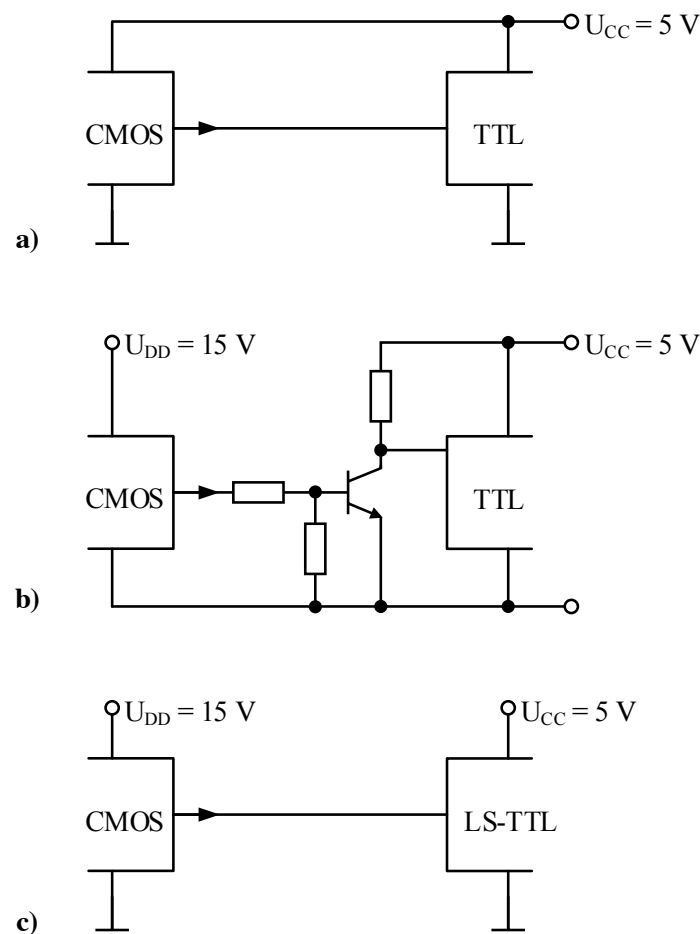


Bild 8.18: Anpassungsschaltungen verschiedener CMOS-Schaltungen mit Versorgungsspannungen von 5 V und 15 V auf TTL- und LS-TTL-Schaltungen.

8.2.3 Anpassungsschaltungen

Müssen innerhalb einer elektronischen Schaltung integrierte Schaltkreise verschiedener Logikfamilien verwendet werden, ist es notwendig, die Art und Weise der Verbindung zu kennen. Die Zusammenschaltung von CMOS-Schaltkreisen mit TTL-Schaltungen bei verschiedenen Versorgungsspannungen zeigt Bild 8.18. Wie Bild 8.18a angibt, können die Ausgänge von CMOS-Bausteinen direkt mit den Eingängen von TTL-Bausteinen verbunden werden, wenn eine einheitliche Versorgungsspannung von $U_{CC} = 5\text{ V}$ vorhanden ist. Werden CMOS-IC's mit einer Versorgungsspannung von 15V eingesetzt, muss zur Ansteuerung von Standard-TTL Bausteinen eine Schaltung nach Bild 8.18b zur Anpassung der Logikpegel aufgebaut werden. Die Inversion des Signals durch die Anpassungsschaltung stellt kein Problem dar, weil Sie durch einen weiteren Inverter auf der CMOS- oder TTL-Seite leicht kompensiert werden kann.

Eine Anpassungsschaltung entfällt in der Regel, wenn anstelle von Standard-TTL,

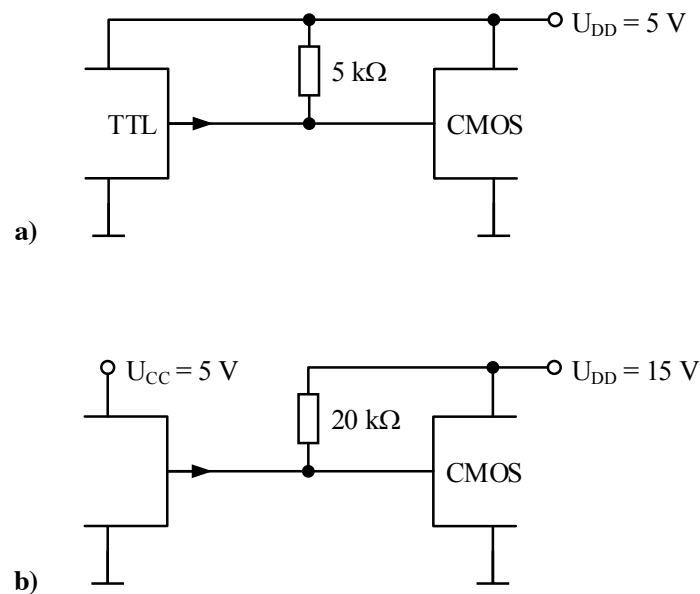


Bild 8.19: Anpassungsschaltungen von TTL- auf CMOS-Schaltungen.

Low-Power-Schottky-TTL-Bausteine verwendet werden (Bild 8.18c), da bei Eingangsspannungen $> 5\text{ V}$ die Schottky-Diode in Sperrrichtung betrieben wird. Bei einer Verbindung des Ausgangs eines Gatters in TTL-Technologie mit Eingängen von CMOS-Schaltkreisen, sollte immer ein sogenannter Pull-Up-Widerstand nach den Bildern 8.19a und 8.19b eingesetzt werden, um den H-Pegel des TTL-Ausgangs anzuheben. Damit wird erreicht, dass am CMOS-Eingang auch eindeutig der H-Pegel anliegt. Auch wird dadurch der Störabstand am Eingang der CMOS-Schaltung vergrößert.

8.2.4 CMOS-Baureihe 4000

Diese CMOS-Baureihe war die erste funktionsfähige Reihe in CMOS-Technologie, die vergleichbare Funktionen zur TTL-Schaltkreisfamilie anbot. Jeder Schaltkreis dieser Familie enthält zusätzliche Ausgangspuffer. Das hat mehrere Vorteile: größere Ausgangstreiberfähigkeit, höherer statischer Störabstand infolge der wesentlich steileren Übertragungskennlinie, höherer dynamischer Störabstand wegen des kleineren Ausgangswiderstandes sowie kürzere Verzögerungs- und Schaltzeiten. Die Schaltungen der CMOS-Baureihe 4000 können im Betriebsspannungsbereich von 3 bis 15V betrieben werden. Dabei beträgt die statische Stromaufnahme typischer Gatter bei 5 V Betriebsspannung für Inverter $1\text{ }\mu\text{A}$, für Flip-Flops $4\text{ }\mu\text{A}$ und für höher integrierte Schaltungen ca. $50\text{ }\mu\text{A}$. Die typischen Eingangsrestströme sind kleiner als $0,3\text{ }\mu\text{A}$. Die typische Verzögerungszeit bei einer kapazitiven Last von 50 pF beträgt bei 5 V 55 ns für ein einfaches NAND-Gatter.

8.2.5 Hochgeschwindigkeits-HC/HCT-Reihe

Seit 1981 werden schnelle CMOS-Logikschaltkreise hergestellt, die etwa die Geschwindigkeit der Low-Power Schottky-TTL-Familie aufweisen. Sie haben den Vorteil, dass sie bedingt durch die geringe Leistungsaufnahme bis in den Frequenzbereich von einigen MHz wesentlich niedrigere Verlustleistung haben. Diese höhere Schaltgeschwindigkeit wurde dadurch erzielt, dass anstelle des bei der 4000er Reihe verwendeten Metall-Gates ein Silizium-Gate Verwendung findet. Die Herstellung der schnellen HC/HCT-Reihen führte in großem Umfang zur Ablösung der bis zu diesem Zeitpunkt dominierenden TTL- bzw. Low-Power Schottky-TTL-Schaltkreise. Es gibt zahlreiche Signal- und Anschlusskompatible Typen. Die HC-Reihe kann mit Ausgangssignalen von TTL-Schaltkreisen angesteuert werden, muss aber am Eingang mit einem 'Pull-up-Widerstand' beschaltet werden, um den Signalpegel geringfügig anzuheben. Die HCT-Reihe unterscheidet sich gegenüber den anderen Baureihen durch die erniedrigte Umschaltsschwelle. Dadurch ist diese Serie voll TTL-kompatibel. Die Hochgeschwindigkeits-CMOS-Reihen HC/HCT haben folgende Vorteile gegenüber den Reihen CMOS 4000 bzw. Low-Power Schottky-TTL:

- Vergleichbare Schaltgeschwindigkeit zu Low-Power Schottky-TTL
- Wesentlich kleinere Leistungsaufnahme als Low-Power Schottky-TTL bei niedrigen und mittleren Schaltfrequenzen
- ca. 10-fache Schaltgeschwindigkeit und ca. 10-fache Ausgangstreiberfähigkeit im Vergleich zur CMOS-Reihe 4000.

Eingangsstufe

Bild 8.20 zeigt die Eingangsstufe und Bild 8.21 die Ausgangsstufe der HCT-Reihe. Durch D_3 und dadurch, dass T_2 einen höheren Stromverstärkungsfaktor als T_1 aufweist, wird die Umschaltsschwelle auf ca. 1,4 V abgesenkt, um dadurch eine volle TTL-Kompatibilität zu erzielen. D_1, D_2, R_1 und R_2 schützen die Eingänge vor elektrostatischen Aufladungen.

Ausgangsstufe

T_6 und T_7 bilden eine Gegentakt-Endstufe. Die Ausgangsstrombelastbarkeit des Schaltkreises liegt je nach Typ in der Größenordnung von 30 bis 60 mA. D_4 und D_5 ergeben sich durch parasitäre pn Übergänge. Ein Standard-HCT-Ausgang kann ≤ 10 Low-Power-Schottky-TTL-Eingänge ansteuern.

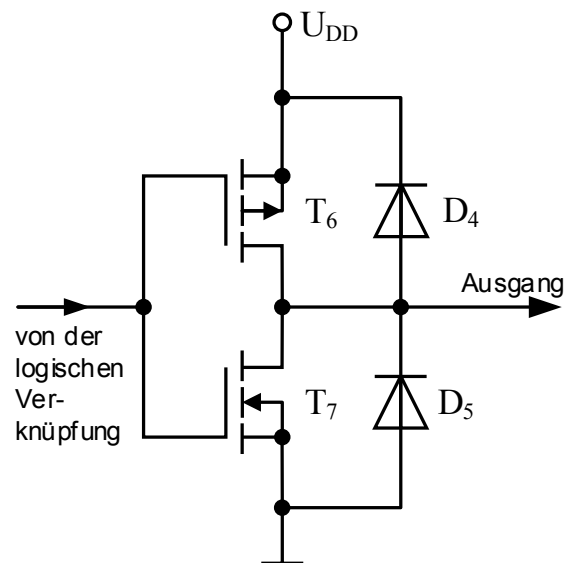


Bild 8.21: Ausgangsstufe der CMOS HCT-Schaltkreisreihe.

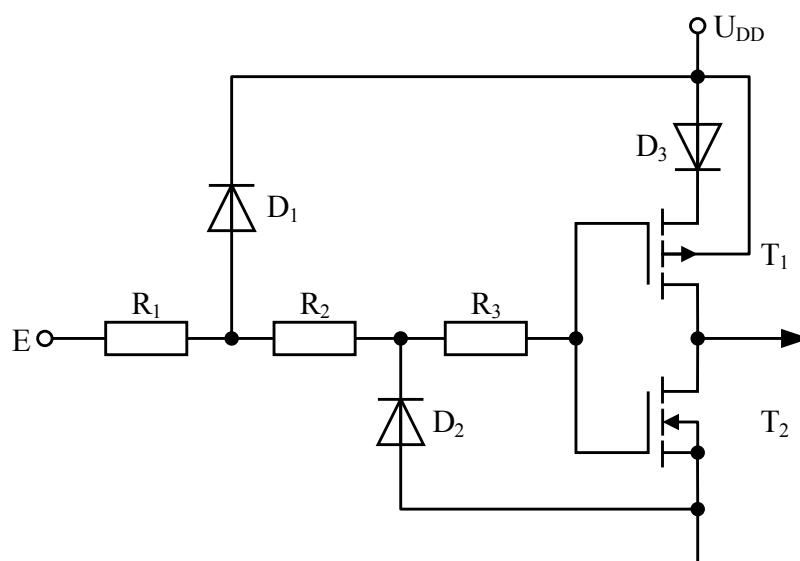


Bild 8.20: Eingangsstufe der CMOS HCT-Schaltkreisreihe.

8.2.6 Advanced CMOS-Reihe AC/ACT

Die wesentlichen Anwendervorteile dieser Schaltkreisreihe sind bei der Beibehaltung der sehr niedrigen CMOS-Verlustleistung eine höhere Schaltgeschwindigkeit gegenüber der HC/HCT-Reihe. So erreichen Toggle-Flip-Flops maximale Frequenzen größer als 100 MHz sowie vergrößerte Ausgangsströme, z.B. von ± 75 mA bei Treiberschaltungen für $75\ \Omega$ Leitungen. Weiterhin haben sie eine deutliche Verringerung der Amplitude von

Ausgangsstörspannungen, die infolge von Resonanzen und Reflexionen im Schaltkreis auftreten können.

8.2.7 Advanced High-Speed-CMOS-Reihe AHC/AHCT

Diese CMOS-Reihe wurde 1995/96 eingeführt und hat die vorteilhaften Eigenschaften der HC/HCT-Reihe, jedoch eine 3-fach höhere Signalverarbeitungsgeschwindigkeit und nur 50 % der statischen Verlustleistung. Diese Reihe wurde optimiert für einen Betriebsspannungsbereich von 2,0 bis 5,5 V und teilweise auch spezifiziert für Betriebsspannungen von 3,3 V. Die Eingangsspannungen dürfen bis zu 5 V betragen. Typische Flip-Flop-Taktfrequenzen können bis zu 170 MHz bei 5 V betragen. Die Ausgangstreiberfähigkeit liegt bei ± 8 mA für 5 V.

Trend

Zusammenfassend wollen wir noch einmal die wichtigsten Vor- und Nachteile von CMOS-Schaltungen aufführen.

Vorteile:

- Extrem niedriger Leistungsverbrauch bei niedrigen Frequenzen (statische Verlustleistung im nW-Bereich)
- Hohe statische und dynamische Störsicherheit ($\approx 0,3$ bis $0,45 \cdot U_E$)
- Extrem großer Betriebsspannungsbereich (2 bis 15 V, Uhrenschaltkreise bis herab zu 0,7 V)
- Weitgehende Temperaturunabhängigkeit der Schaltschwelle.

Nachteile:

- Relativ großer Flächenbedarf (etwa 3-fach gegenüber NMOS-Technik)
- Mehr Herstellungsschritte als bei NMOS-Schaltungen.

Die hohe Störsicherheit, die einfache Betriebsspannungsversorgung (unstabilisiert, sehr geringer Leistungsbedarf, Betriebsspannung von TTL-Systemen verwendbar) und die geringe Leistungsaufnahme bei niedrigen Schaltfrequenzen sind bei zahlreichen Einsatzfällen erhebliche Vorteile gegenüber TTL-Schaltungen. CMOS-Schaltungen sind daher seit Jahren (etwa ab Mitte der 80er Jahre) die wichtigste Schaltungstechnologie der digitalen Schaltungstechnik, insbesondere für höchstintegrierte Schaltkreise. Allerdings wird bei der Groß- und Höchstintegration meist nicht die „klassische“ CMOS-Technik, sondern eine Kombination von CMOS- und NMOS-Technik verwendet.

1. „Klassische“ CMOS-Technik je 50 % p- und n-Kanal MOS, vorwiegend bei Schaltkreisen mittleren Integrationsgrades (μ Prozessoren, SRAM, Logik).

Vorteil: Volle Ausnutzung der Leistungsfähigkeit von CMOS-Strukturen.

2. Verhältnis p– zu n–Kanal MOS–Technik von 40:60 bzw. 20:80, typisch für die meisten 8–bit CMOS–Mikroprozessoren, Kompromiss zwischen vertretbarer Vergrößerung der Chipfläche und verbesserter CMOS–Leistungsfähigkeit.
3. Sehr geringes Verhältnis von p– zu n–Kanal MOS–Technik, d.h. geringer CMOS–Anteil in NMOS–Strukturen, um die Leistungsfähigkeit gegenüber reinen NMOS–Strukturen bei kaum vergrößerter Chipfläche zu verbessern. Beispiel: DRAM.

8.2.8 BiCMOS–Schaltungen

Neben vielen Vorteilen haben CMOS–Schaltkreise den Nachteil, dass ihre Ausgangstreiberfähigkeit kleiner ist als die von Endstufen mit bipolaren Transistoren. Das ist z.B. bei Bustreiberschaltkreisen für sehr hohe Schaltfrequenzen ein Nachteil (Treiben von größeren kapazitiven Lasten). Durch die Kombination der CMOS–Technologie mit bipolaren Transistoren lassen sich die günstigen Eigenschaften der CMOS–Technik mit denen der Bipolartechnik verbinden. Bipolare Transistoren können infolge ihrer hohen Steilheit bei Ansteuerung mit relativ kleiner Basisspannung große Ausgangsströme mit sehr steilen Flanken liefern, so dass Ausgangskapazitäten besonders schnell umgeladen werden. In den Eingangs– und den Zwischenstufen von Gattern sind CMOS–Schaltungen wegen ihrer sehr geringen Leistungsaufnahme überlegen.

Die Firma Texas Instruments brachte 1987 die erste BiCMOS–Reihe auf den Markt. In BiCMOS–Schaltungen werden neben CMOS–Gattern auch npn Transistoren, vorzugsweise in Endstufen, verwendet. Da zusätzliche technologische Schritte gegenüber CMOS–Schaltungen erforderlich sind und der Chipflächenbedarf gegenüber CMOS bis zu einem Faktor 2 größer ist, steigen die Herstellungskosten. BiCMOS–Schaltkreise verwenden TTL–Pegel an den eigenen Ausgängen. Sie sind mit TTL–Familien direkt zusammenschaltbar. Gegenüber Bipolarschaltkreisen haben sie vor allem im hochohmigen Zustand von Tristate–Ausgangsstufen einen deutlich geringeren Leistungsbedarf ($< 10\%$). Mit steigender Schaltfrequenz nimmt ihre Leistungsaufnahme nur wenig zu, vergleichbar mit Bipolarschaltkreisen.

Eingangsschaltung

Bild 8.22 zeigt die vereinfachte Eingangsschaltung der eines BiCMOS–Gatters. Um am Eingang Kompatibilität mit TTL–Pegeln zu erhalten, muss die Umschaltswelle des MOS–Inverters T_2, T_3 etwa 1,5 V betragen. Die Betriebsspannung des Inverters wird daher mittels D_1 und T_1 um etwa 1,5 V verringert. Das hat zusätzlich den Vorteil, dass der beim Umschalten des Inverters auftretende Stromspike zwischen U_{DG} und Masse verkürzt und verkleinert wird. Für $U_E = L$ wird durch einen Rückkopplungsweig, in den T_4 einbezogen ist, die Sourcespannung von T_2 auf die volle Betriebsspannung U_{CC} erhöht, damit die nachfolgenden Stufen sicher angesteuert werden. Zusätzlich erzeugt diese Rückkoppelungsschaltung eine Eingangsschalthysterese von 100 bis 200 mV, wodurch der Störabstand vergrößert wird.

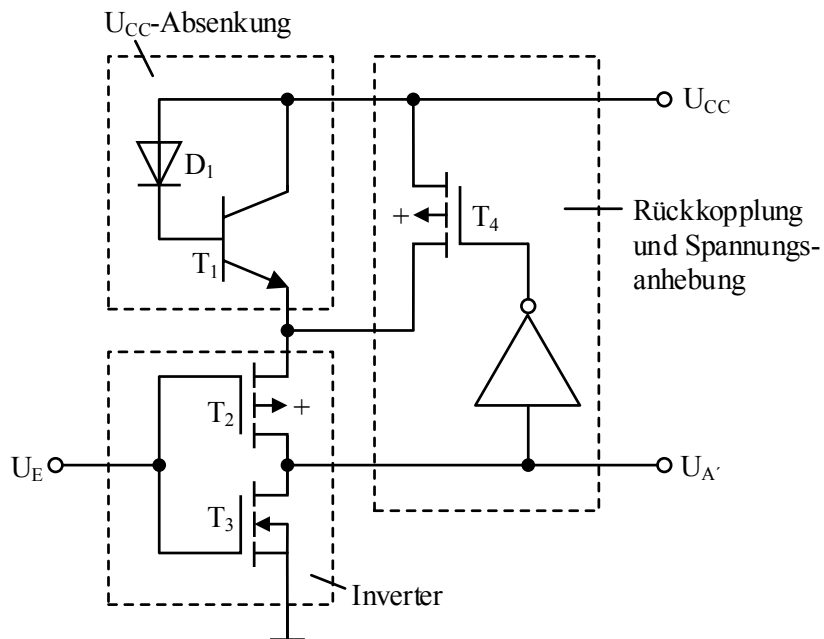


Bild 8.22: Eingangsstufe eines BiCMOS-Gatters.

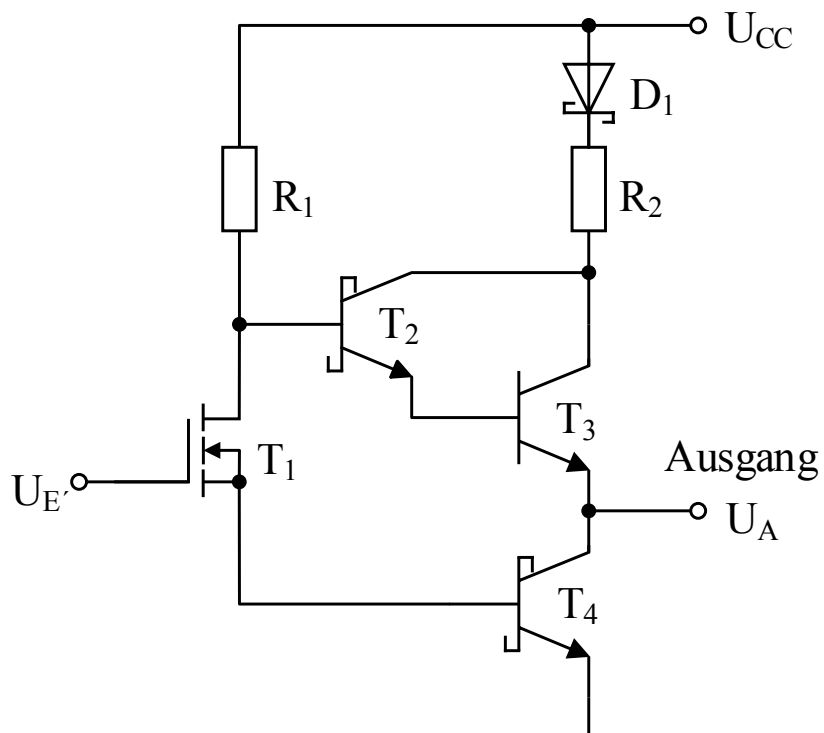


Bild 8.23: Ausgangsstufe eines BiCMOS-Gatters.

Ausgangsschaltung

Bild 8.23 zeigt vereinfacht eine BiCMOS–Ausgangsstufe. Wenn der Stromschalter T_1 leitet, fließt Strom über R_1 und T_1 in die Basis von T_4 . Folglich gilt: $U_A = U_{AL}$. Zur gleichen Zeit sperren T_2 und T_3 . Wenn U_E von H auf L schaltet, sperren T_1 und T_4 und das Darlington–Paar T_2, T_3 wird über den durch R_1 fließenden Basisstrom eingeschaltet. R_2 begrenzt den Ausgangsstrom bei $U_A = U_{AH}$. 1997 erschien die mit einer Betriebsspannung von 3,3 V betriebene ALB-(Advanced Low-Voltage BiCMOS) Familie.

In Tabelle 8.4 sind die Schaltkreisfamilien mit den wichtigsten Merkmalen aufgeführt.

Schaltkreisfamilie (Baureihen)	Typische Grundschaltung eines Gatters (vereinfacht)	Vor- und Nachteile	Bemerkungen zur Funktion	Bemerkungen zum Einsatz in der Informationstechnologie
n-Kanal MOS		+ einfacher Aufbau (nur MOS-FET) + sehr gut für IS geeignet + nahezu leistungslose Ansteuerung - etwas niedrigere Schaltgeschwindigkeit als TTL	• direkte Kopplung der Stufen	• größerer Störabstand als TTL, vorteilhafter sind CMOS-Schaltkreise • veraltet, wird durch CMOS abgelöst
CMOS (Komplementär-MOS)		+ hoher statischer Störabstand ($0,45 \cdot U_S$) + großer Betriebsspannungsbereich; nicht TTL-kompatibel + geringer Ruheleistungsverbrauch ($P_V \leq 100 \text{ nW}$, Ruhestrom $\approx 10 \text{ nA}$) + Schaltgeschwindigkeit fast vergleichbar mit TTL - teurer als n-Kanal MOS	• direkte Kopplung der Stufen • merklicher Leistungsverbrauch nur beim Umschaltvorgang • Reihe 4000: $U_S = 3..15 \text{ V}$ • Reihe HC: $U_S = 2..6 \text{ V}$ • Reihe HCT: $U_S = 4,5..5 \text{ V}$	• wegen des hohen Störabstandes, der einfachen Betriebspannungsvorsorgung (unstabilisiert) und des sehr geringen Leistungsverbrauchs sehr gut für den Einsatz in der IT geeignet • wurde zur Standard-Schaltkreisfamilie
BiCMOS		+ geringer Verbrauch + sehr hohe Ausgangstreiberfähigkeit + kurze Schaltzeit bei kapazitiver Last - erhöhte Herstellungskosten	• Verbindung der Vorteile von CMOS und Bipolar-technik	• besonders als Bustreiber für sehr hohe Schaltfrequenzen

Tabelle 8.4: Übersicht der bipolaren Logikschaltkreise

8.2.9 Elektronische Schalter (Transferrgatter, Transmission Gate)

In der Schaltungstechnik werden neben den logischen Gattern häufig auch schnelle elektronische Schalter benötigt, die immer dann eingesetzt werden, wenn ein mechanischer Schalter oder ein Relais zu langsam sind, oder wenn ein Schalter in einer integrierten Schaltung benötigt wird.

Bild 8.24a zeigt den Aufbau eines heute üblichen elektronischen Schalters in der CMOS-Technik. Er besteht aus einem n- und einem p-Kanal Transistor, die parallel geschaltet sind und mit invertierten Steuersignalen U_{St} und $\overline{U_{St}}$ betrieben werden. Ist $U_{St} = 0$ V (Low-Pegel) und $\overline{U_{St}} = U_{DD}$ (High-Pegel), sperren beide Transistoren. Dies entspricht einem geöffneten mechanischen Schalter.

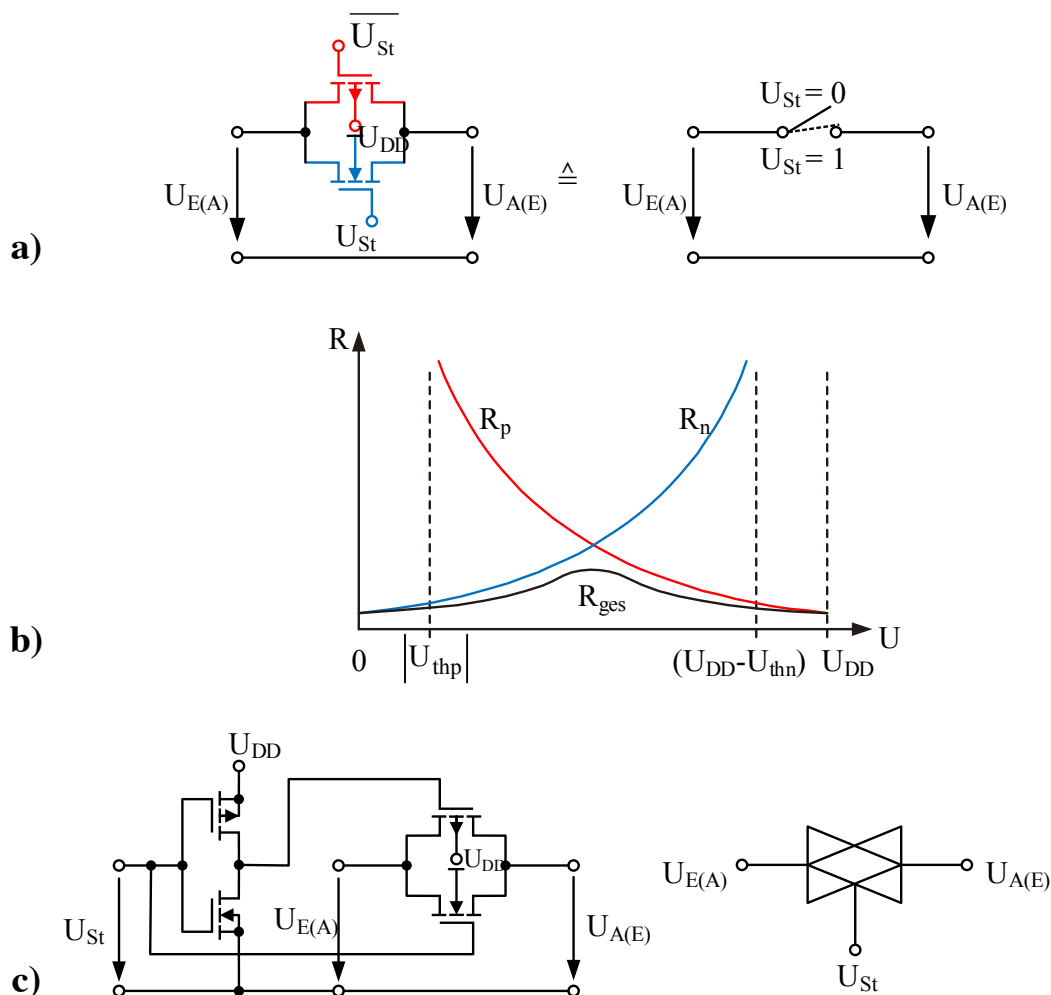


Bild 8.24: a) Elektronischer Schalter in CMOS-Technik, b) Widerstand des Schalters für $U_{St} = 1$, c) komplette Schaltung mit Inverter zum Umschalten des Transferrgatters und logisches Symbol.

Ist $U_{St} = U_{DD}$ (High-Pegel) und $\overline{U_{St}} = 0$ V (Low-Pegel), leiten beide Transistoren. Ein mechanischer Schalter wäre jetzt geschlossen. Damit wird erreicht, dass beide Transistoren gleichzeitig aus- bzw. eingeschaltet sind. In Bild 8.24b sind die Widerstände R_n und R_p der beiden Transistoren und der Gesamtwiderstand R_{ges} des elektronischen Schalters im eingeschalteten Zustand dargestellt. Liegen kleine positive Eingangsspannungen U_E an, so ist der n-Kanal Transistor niederohmig. Bei Eingangsspannungen $|U_{th,p}| < U_E < (U_{DD} - U_{th,n})$ ist der Gesamtwiderstand etwa konstant und bei Spannungen $U_E > (U_{DD} - U_{th,n})$ ist der p-Kanal Transistor niederohmig.

Damit ist es möglich, sowohl analoge Spannungswerte wie auch digitale Pegel von der einen Seite zur anderen Seite des elektronischen Schalters zu übertragen, wie dies mit jedem mechanischen Schalter möglich ist. Werden Umschalter benötigt, so müssen lediglich zwei Transmission Gates parallel geschaltet und auf einer Seite miteinander verbunden werden.

9 Sequentielle Logik

Lernziele

- Kennenlernen von Grundelementen für die sequentielle Logik
- Klassifizierung der verschiedenen Flip-Flop Schaltungen
- Schaltkreisfamilien in CMOS-Technik
- Verstehen der Funktion von verschiedenen Flip-Flop-Familien:
 - zustandsgesteuerten Flip-Flops
 - taktzustandsgesteuerten Flip-Flops
 - taktflankengesteuerten Flip-Flops

9.1 Klassifizierung

Allen Flip-Flop ist gemeinsam, dass sie zwei stabile Zustände haben. Sie unterscheiden sich aber wesentlich durch die Bedingungen, unter denen sie zwischen den beiden Zuständen hin und zurück schalten. Es lassen sich nach Bild 9.1 zwei große Gruppen bilden. Eine Gruppe umfasst alle Flip-Flop, die zustandsgesteuert sind, während die zweite Gruppe alle taktgesteuerten (getakteten) Flip-Flop-Arten erfasst. Unter Taktsteuerung versteht man folgenden Vorgang: die Übernahme der Information von den vorbereitenden Eingängen erfolgt nur bei Vorhandensein eines zusätzlichen Taktsignals.

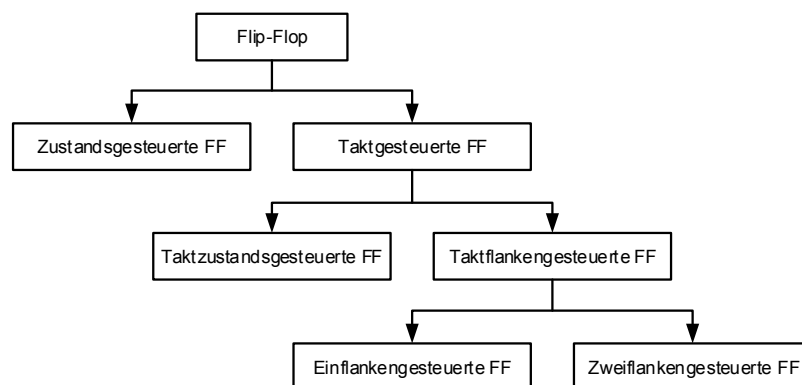


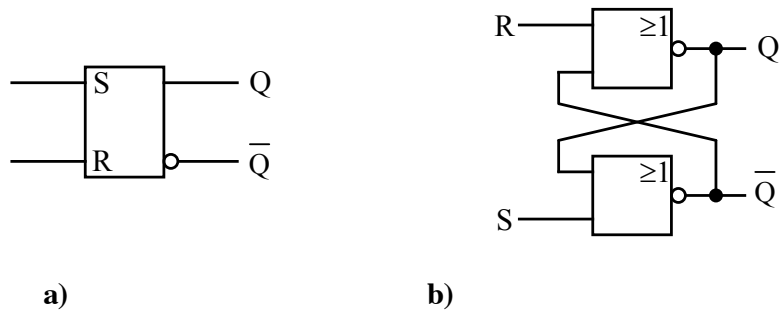
Bild 9.1: Übersicht über die Flip-Flop-Gruppen.

Die Gruppe der taktgesteuerten Flip-Flop kann wieder unterteilt werden in zwei Gruppen: 1. die Gruppe der taktzustandsgesteuerten und 2. die Gruppe der taktflankengesteuerten Flip-Flop. Beim einfachen taktzustandsgesteuerten Flip-Flop wird die an den Vorbereitungseingängen liegende Information während der gesamten Dauer des Taktimpulses in das Flip-Flop übernommen. Somit wirken sich Signaländerungen während der Taktphase sofort auf den Ausgang aus. Deshalb werden sie als Auffang-Flip-Flop (Latch) benutzt und sind für Teiler und Zähler nicht geeignet. Werden zwei taktzustandsgesteuerte Flip-Flops in Reihe geschaltet (Master-Slave), wirkt sich eine Signaländerung an den Vorbereitungseingängen zunächst nur auf den Master aus. Erst wenn der Taktzustand sich ändert (z.B. von H nach L) wird der im Master gespeicherte Zustand über den Slave-Flip-Flop an den Ausgang übergeben. Bei flankengesteuerten Flip-Flop wird die Information mit der aktiven Taktflanke übernommen. Während der Zeit zwischen zwei aktiven Taktflanken ist das Flip-Flop unempfindlich gegenüber Signaländerungen an den Vorbereitungseingängen. Die Flanken des Taktimpulses müssen sehr steil sein, damit keine falschen Informationen übernommen werden. Eine weitere Möglichkeit der Flankensteuerung besteht auch hier in der Zusammenschaltung von zwei Flip-Flops zu einem Master-Slave Flip-Flop. Dabei werden beide Flanken des Taktimpulses ausgenutzt. Deshalb nennt man diese Steuerung Zweiflankensteuerung. Der Eingang des ersten Flip-Flops spricht auf die aufsteigende Taktflanke an, während das zweite Flip-Flop erst auf die abfallende Flanke reagiert.

9.1.1 Zustandsgesteuerte Flip-Flop

Ein RS-Flip-Flop, dessen logisches Schaltbild in Bild 9.2a gezeigt ist, kann als einfachste Form einer Speicherschaltung betrachtet werden. Bild 9.2b zeigt den Aufbau eines RS-FF mit NOR-Gattern. Kennzeichnend für Flip-Flops aus Gattern ist die überkreuzte Rückkopplung der Ausgänge auf die Eingänge der Gatter. Dadurch wird erreicht, dass die Ausgangszustände erhalten bleiben, auch wenn das entsprechende Eingangssignal nicht mehr anliegt. Die Wahrheitstabelle für das RS-Flip-Flop ist in Bild 9.2c angegeben. Sind die Eingänge $S = R = 0$, so ist das Flip-Flop inaktiv, d.h. der gespeicherte Zustand bleibt erhalten. Eine Kombination der Eingangssignale $S = 1$, $R = 0$ bewirkt, dass der Ausgang $Q = 1$ wird, während mit $S = 0$ und $R = 1$ der Ausgang $Q = 0$ wird. Der Zustand $S = R = 1$ erzeugt an den Ausgängen beider Gatter den Zustand $Q = \overline{Q} = 0$ (schraffierte Bereiche in Bild 9.2d). Ein Übergang auf den Eingangszustand $S = R = 0$ würde am Ausgang zu einem nicht definierten Zustand $Q = 0$ oder $Q = 1$ führen, abhängig von geringfügigen Unterschieden der Gatterlaufzeiten der beiden NOR-Gatter. Deshalb muss darauf geachtet werden, dass die beiden Eingänge nicht gleichzeitig 1 werden können.

In Bild 9.3a ist das RS-Flip-Flop mit NOR-Gattern in einer anderen Form gezeichnet. Eine Realisierung eines solchen Flip-Flops in NMOS-Technik erkennt man in Bild 9.3b. Als Lasttransistoren werden hier n-Kanal MOSFETS vom Verarmungstyp eingesetzt. Bild 9.3c zeigt ein RS-Flip-Flop mit zwei NOR-Gattern in CMOS-Technik.



c)

S	R	Q	$\bar{Q}$
0	0	Q_{-1}	$\bar{Q}_{-1}$
0	1	0	1
1	0	1	0
1	1	0	0

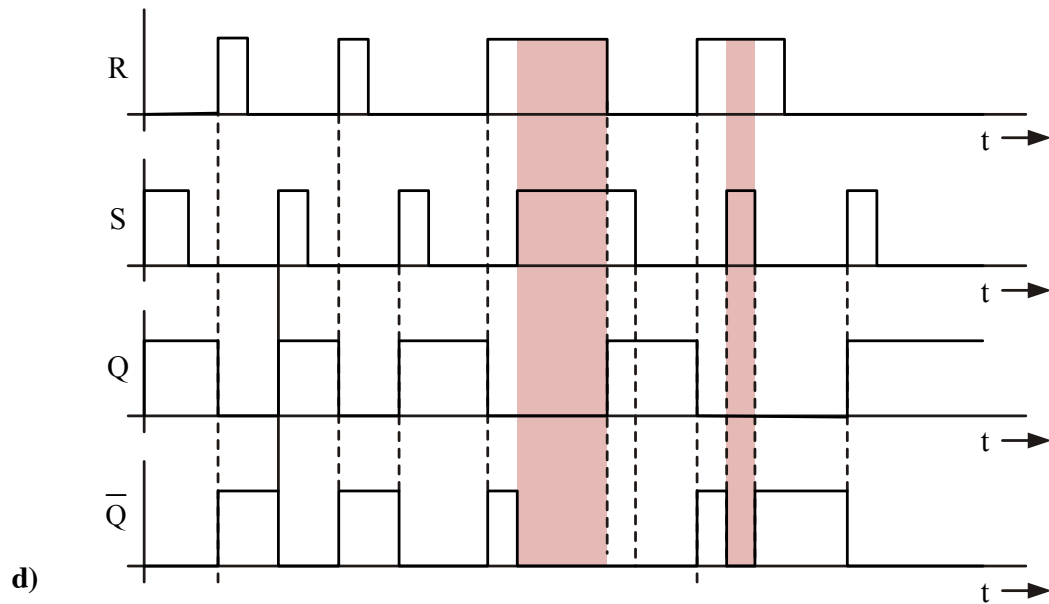


Bild 9.2: RS - Flip-Flop: a) logisches Symbol, b) Schaltbild mit NOR - Gattern und c) Wahrheitstabelle, d) Signalverläufe an den Ein- und Ausgängen.

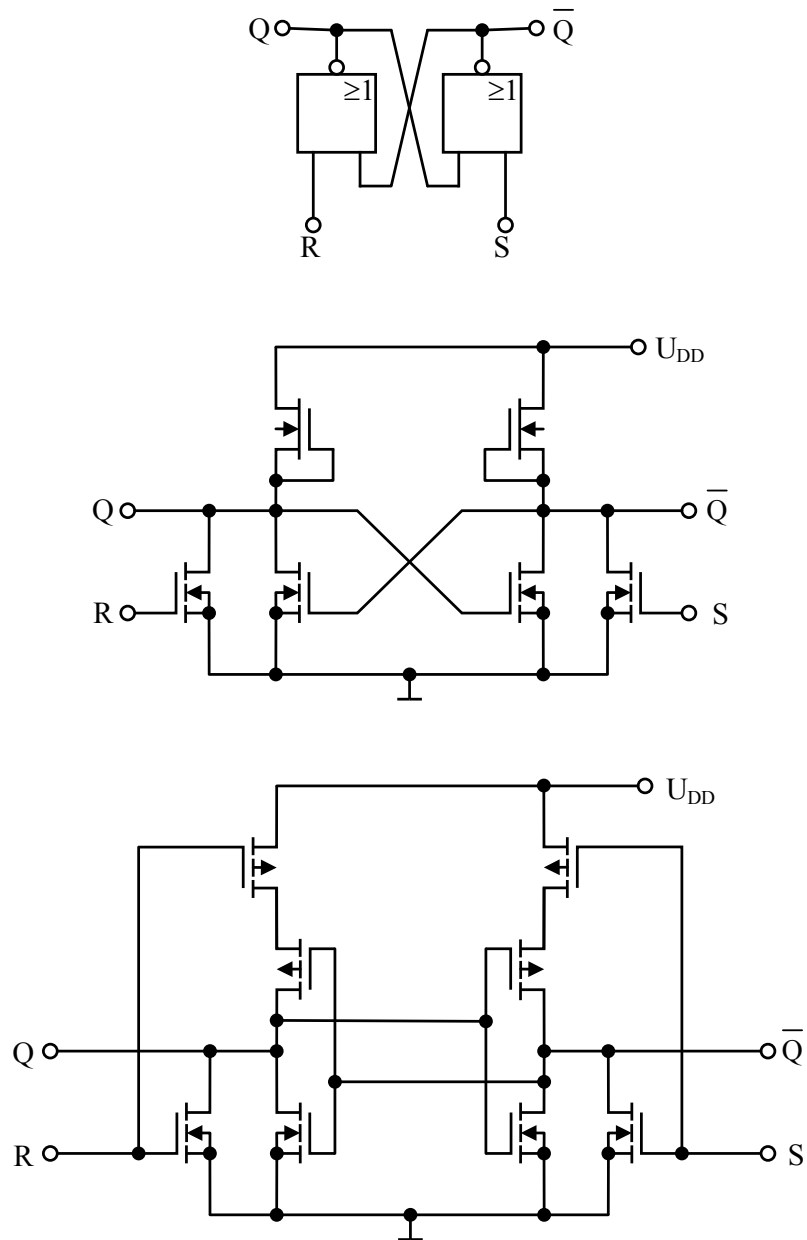


Bild 9.3: RS-Flip-Flop: a) Schaltbild mit NOR-Gattern, b) Schaltung mit n-Kanal MOSFETs vom Anreicherungs- und Verarmungstyp und c) CMOS-Schaltung.

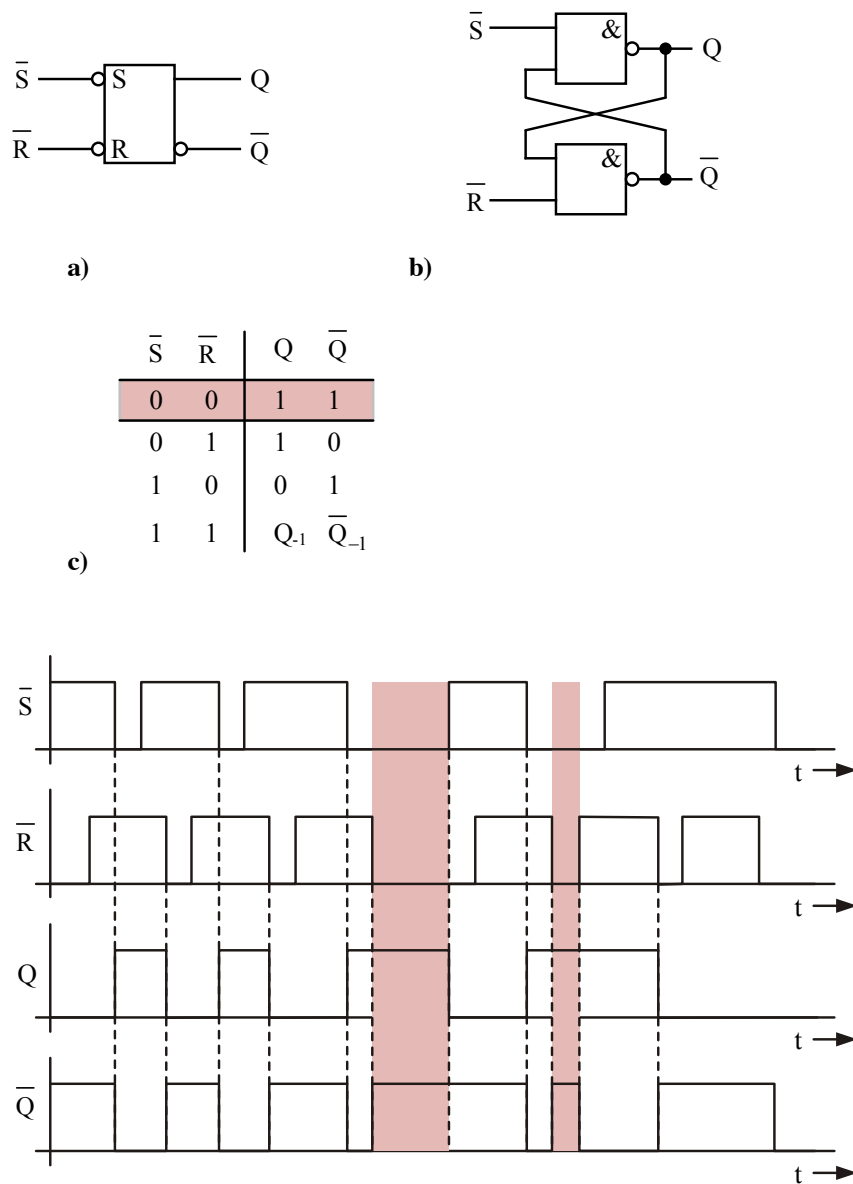


Bild 9.4: $\overline{RS}$ -Flip-Flop: a) logisches Symbol, b) Schaltbild mit NAND-Gattern und c) Wahrheitstabelle, d) Signalverläufe am Ein- und Ausgang.

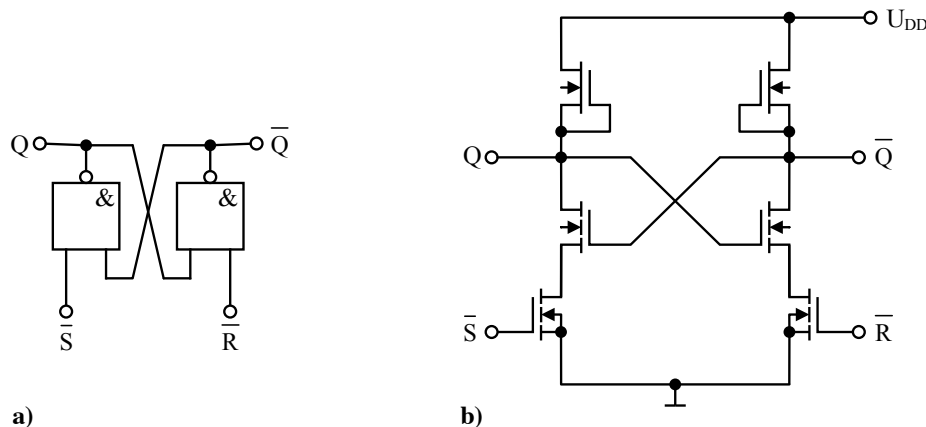


Bild 9.5: $\overline{RS}$ -Flip-Flop: a) Schaltbild mit NAND-Gattern und b) Schaltung mit n-Kanal MOSFETs vom Anreicherungs- und Verarmungstyp.

Verwendet man anstelle der NOR-Gatter NAND-Gatter, erhält man ein $\overline{RS}$ -Flip-Flop. Bild 9.4 zeigt a) das logische Schaltbild, b) den Aufbau mit NAND-Gattern, c) die zugehörige Wahrheitstabelle und d) ein Beispiel von möglichen Signalen an den Eingängen und Ausgängen des Flipflops. Bei diesem Flip-Flop sind die aktiven Eingangszustände die Low-Pegel. Beim $\overline{RS}$ -Flip-Flop ergibt folglich ein Eingangszustand $\overline{R} = \overline{S} = 0$ an den beiden Ausgängen Q und $\overline{Q}$ eine 1, so dass beim gleichzeitigen Übergang der Eingänge nach $\overline{R} = \overline{S} = 1$ der Ausgangszustand nicht definiert ist. Bei diesem Flip-Flop muss also darauf geachtet werden, dass an den beiden Eingängen nicht gleichzeitig ein Low-Pegel anliegt oder ein Übergang von 0 nach 1 nicht gleichzeitig an beiden Eingängen erfolgt. Auch hier sind die Bereiche mit den Eingangszuständen $\overline{R} = \overline{S} = 0$ schraffiert. Bild 9.5 zeigt nochmals das $\overline{RS}$ -Flip-Flop mit dem Aufbau mit NAND-Gattern und eine entsprechende Auslegung der Schaltung mit n-Kanal Feldefekttransistoren. Die Lasttransistoren sind auch hier selbstleitend.

9.1.2 Taktzustandsgesteuerte Flip-Flop

Taktzustandsgesteuertes RS-Flip-Flop

Erweitert man das, RS-Flip-Flop nach Bild 9.2b durch Vorschalten zweier AND-Gatter an die ein gemeinsamer Takt angeschlossen wird, erhält man ein taktzustandsgesteuertes RS-Flip-Flop. Bild 9.6a zeigt das logische Schaltbild und Bild 9.6b den beschriebenen Schaltungsaufbau mit Gattern. Eine Realisierung mit n-Kanal Feldefekttransistoren ist in Bild 9.6c gezeigt. Die zugehörige Wahrheitstabelle nach Bild 9.6d zeigt, dass es auch bei diesem Flip-Flop einen unerlaubten Zustand $R = S = C = 1$ gibt, denn hier sind beide Ausgänge 0. Deshalb muss auch bei diesem Flip-Flop darauf geachtet werden, dass spätestens beim Übergang des Taktsignals C von 1 nach 0 der Zustand $R = S = 1$ vermieden wird, da sonst der Zustand der Ausgänge Q und $\overline{Q}$ nicht definiert ist. In Bild 9.6e ist dieser Zustand $R = S = C = 1$ schraffiert eingezeichnet.

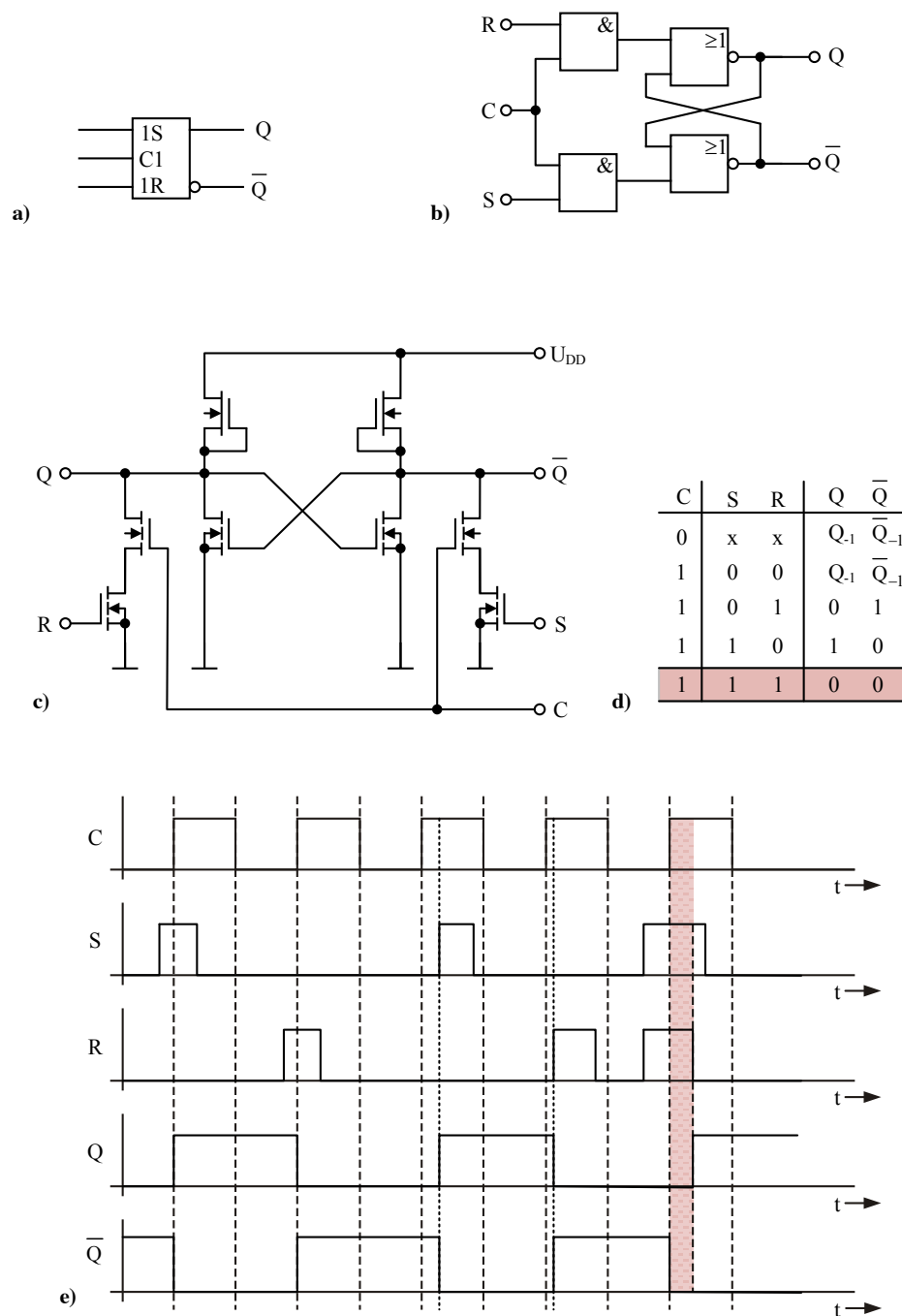


Bild 9.6: Taktzustandsgesteuertes RS-Flip-Flop: a) logisches Symbol, b) Schaltbild mit Gattern, c) Schaltung mit n-Kanal MOSFETs, d) Wahrheitstabelle und e) Signalverläufe.

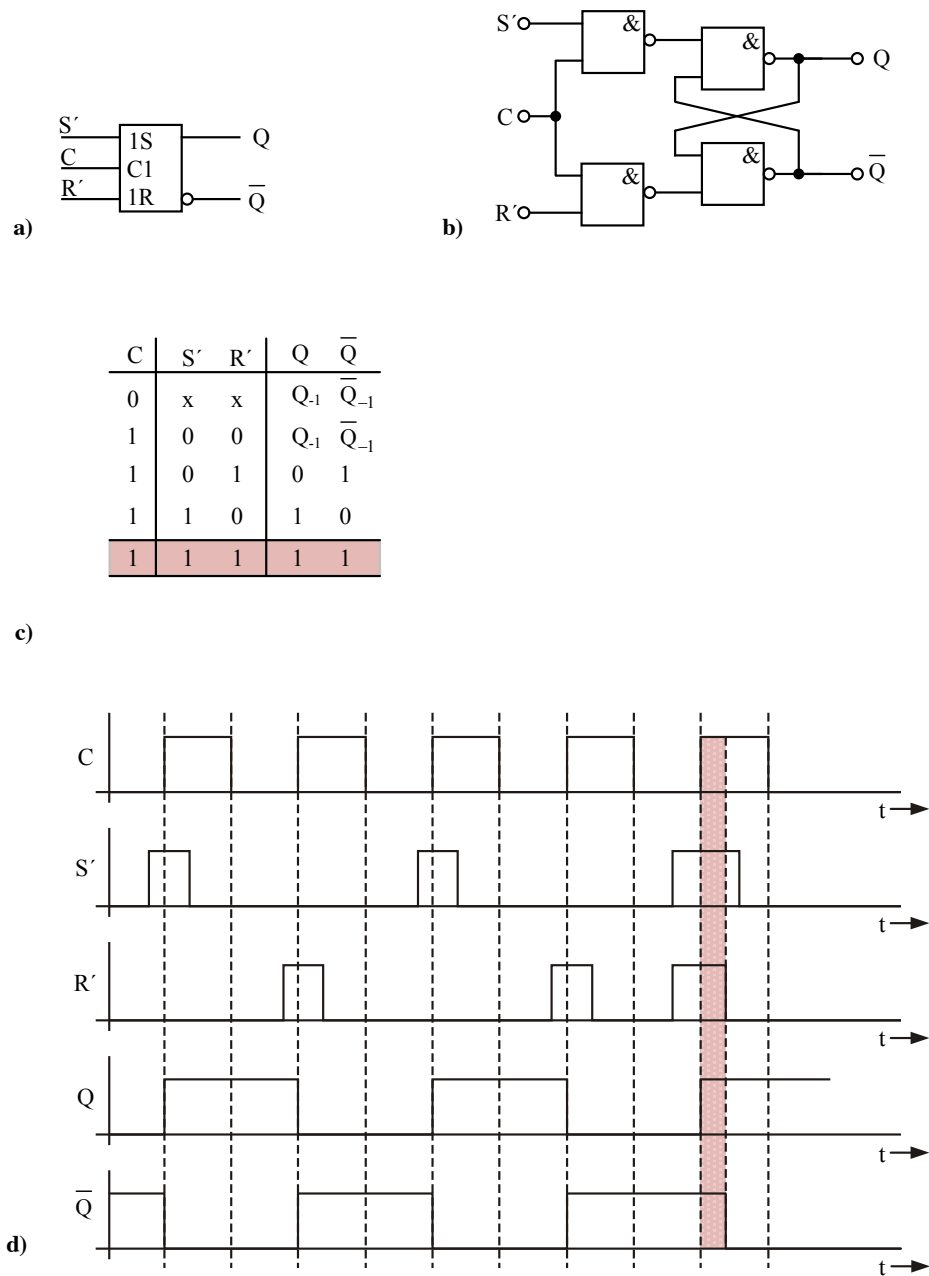


Bild 9.7: Taktzustandsgesteuertes RS-Flip-Flop: a) logisches Symbol, b) Schaltbild mit NAND-Gattern, c) Wahrheitstabelle und d) Signalverläufe.

Vor dem Ende des Taktes ist wieder ein erlaubter Zustand $R = 0$ und $S = 1$ vorhanden, so dass beim Übergang von $C = 1 \rightarrow C = 0$ ein definierter Ausgangszustand vorhanden ist.

Ein weiteres taktzustandsgesteuertes RS-Flip-Flop zeigt Bild 9.7, jedoch mit einem $\overline{RS}$ -Flip-Flop als Speicherschaltung.

Taktzustandsgesteuertes D-Flip-Flop

Ein taktzustandsgesteuertes Flip-Flop bei dem es keinen unerlaubten Eingangszustand gibt, ist das D-Flip-Flop. Dieses Flip-Flop nennt man auch transparentes Flip-Flop, da während der Dauer des H-Pegels des Taktes eine Veränderung der Eingangsinformation D sofort am Ausgang erscheint. Nimmt der Takt den L-Pegel an, bleibt die zuletzt anliegende Information gespeichert. Bild 9.8 zeigt eine mögliche Realisierung eines D-Flip-Flops mit NAND-Gattern. Betrachtet man den D-Eingang als "Set"-Signal, liegt am zweiten Eingangsgatter das invertierte Signal als "Reset" an, wenn gleichzeitig $C = 1$ ist. Ist $C = 0$ ist der D-Eingang inaktiv. In Bild 9.8c die zugehörige Wahrheitstabelle gezeigt. Die Signalverläufe in Bild 9.8d zeigen die beschriebene 'Transparenz' des D-Flip-Flops. Diese Bereiche sind unterlegt eingezeichnet.

Taktzustandsgesteuerte Master-Slave Flip-Flops

Im Weiteren soll ein taktzustandsgesteuertes RS-Master-Slave Flip-Flop nach Bild 9.9 ausführlich erklärt werden. Während des H-Zustands des Taktes wird der Eingangszustand vom ersten Flip-Flop aufgenommen und dort gespeichert, während das zweite Flip-Flop noch die Information des vorigen Zustands speichert. Ändert sich der Pegel des Taktes nach Low, wird die Ausgangsinformation des ersten Flip-Flops in das zweite Flip-Flop aufgenommen und an den Ausgang weitergeleitet. Das erste Flip-Flop wird inaktiv. Die Eingangsinformation erscheint also um die Dauer des H-Zustands des Taktsignals verzögert am Ausgang. Auch bei diesem RS-Flip-Flop muss darauf geachtet werden, dass der Eingangszustand $R = S = C = 1$ vermieden wird (schraffierter Bereich in Bild 9.9c, da sonst beim Übergang des Taktsignals von High nach Low an den Ausgängen $Q2$ und $\overline{Q2}$ ein undefinierter Zustand auftreten kann.

Dieser Nachteil wird vermieden, wenn man anstelle des RS-Flip-Flops ein taktzustandsgesteuertes JK-Master-Slave Flip-Flop verwendet, dessen Aufbau in Bild 9.10 gezeigt ist. (J = Jump, K = Kill). Der entscheidende Unterschied zum RS-Flip-Flop liegt in der Rückführung der Ausgangssignale $Q2$ und $\overline{Q2}$ und deren Verknüpfung mit den Eingängen K und J. Durch die so entstandene Verknüpfung von C, J und $\overline{Q2}$ bzw. C, K und $Q2$ kann bei $J = K = 1$ immer nur ein Zweig des Eingangs-Flip-Flops gesetzt werden. Man spricht in diesem Fall auch von einem "Toggle-Flip-Flop". Die Signalverläufe in Bild 9.10d zeigen die Zustände an den Eingängen und an den Ausgängen des Master- ($Q1$) und des Slave-Flip-Flops ($Q2$). Die schraffierten Bereiche markieren die

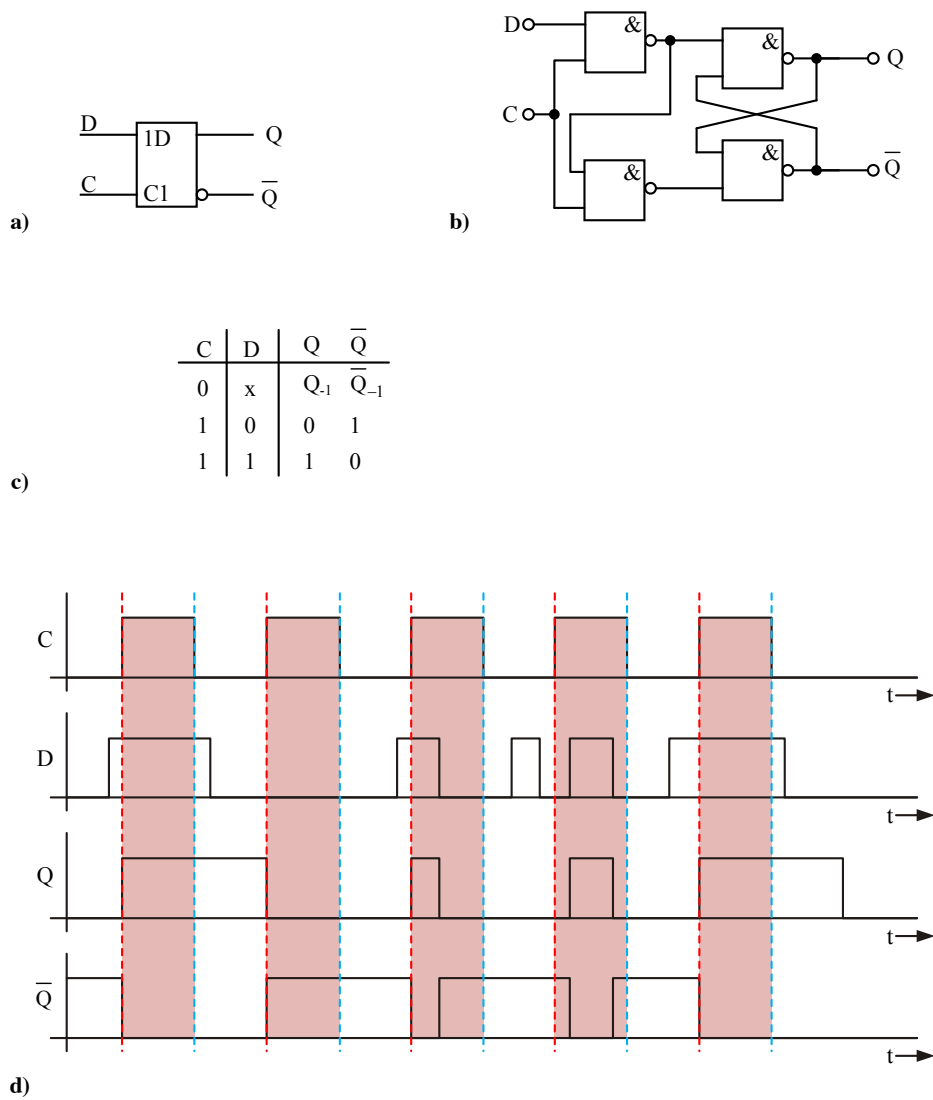


Bild 9.8: Taktzustandgesteuertes D-Flip-Flop: a) logisches Symbol, b) Schaltbild mit NAND-Gattern, c) Wahrheitstabelle und d) Signalverläufe.

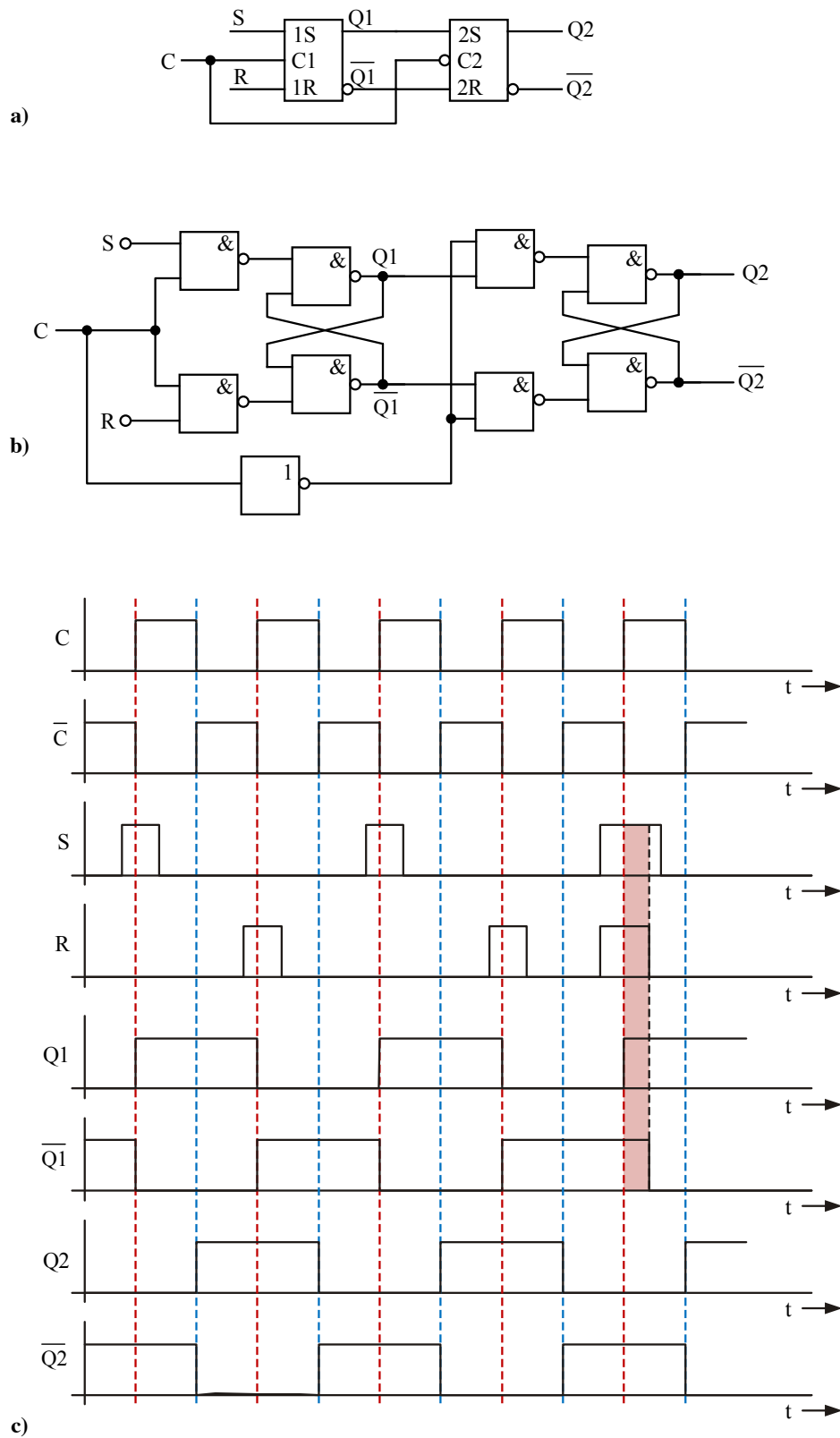


Bild 9.9: Taktzustandgesteuertes RS - Master - Slave - Flip-Flop: a) Schaltbild mit logischen Symbolen b) Schaltbild mit logischen Gattern und c) Signalverläufe.

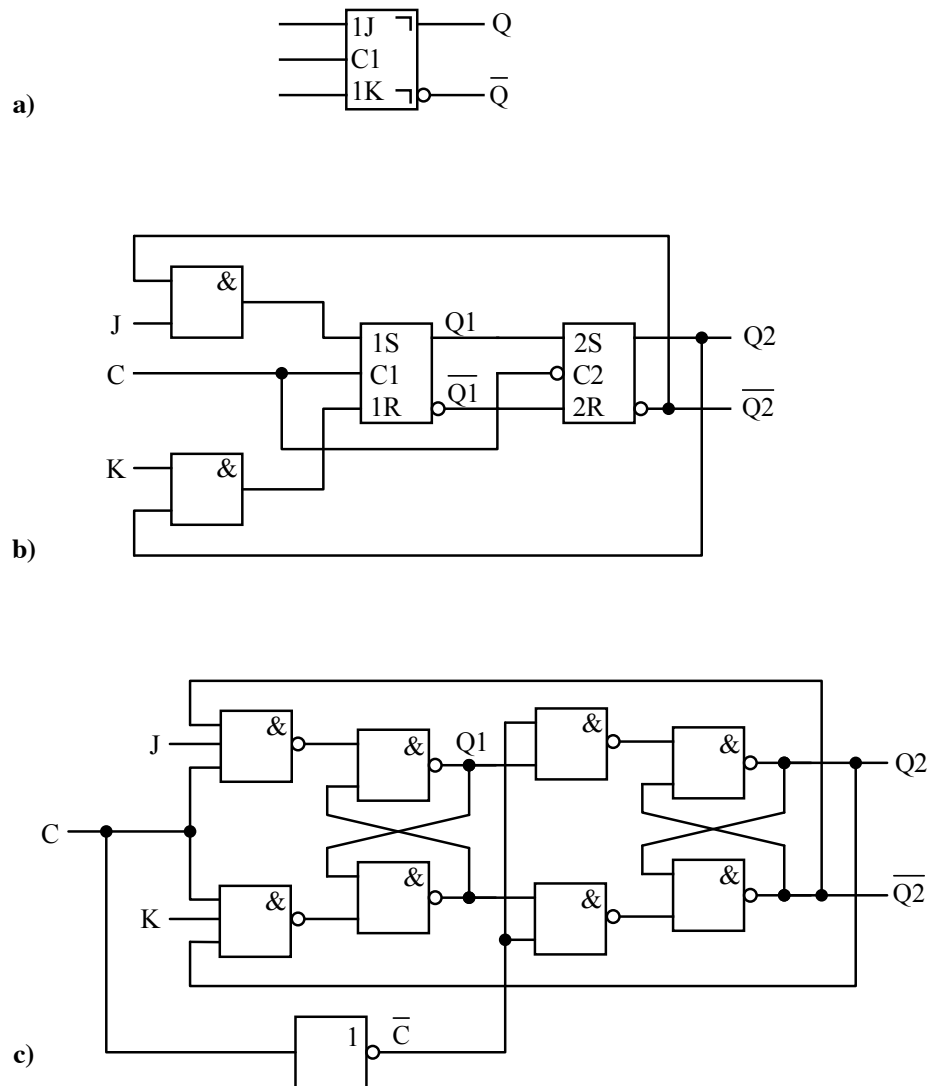


Bild 9.10: Taktzustandgesteuertes JK-Master-Slave-Flip-Flop mit Verzögerung: a) logisches Symbol, b) Schaltbild mit taktzustandgesteuerten RS-Flip-Flops und Gattern, c) Schaltbild mit logischen Gattern.

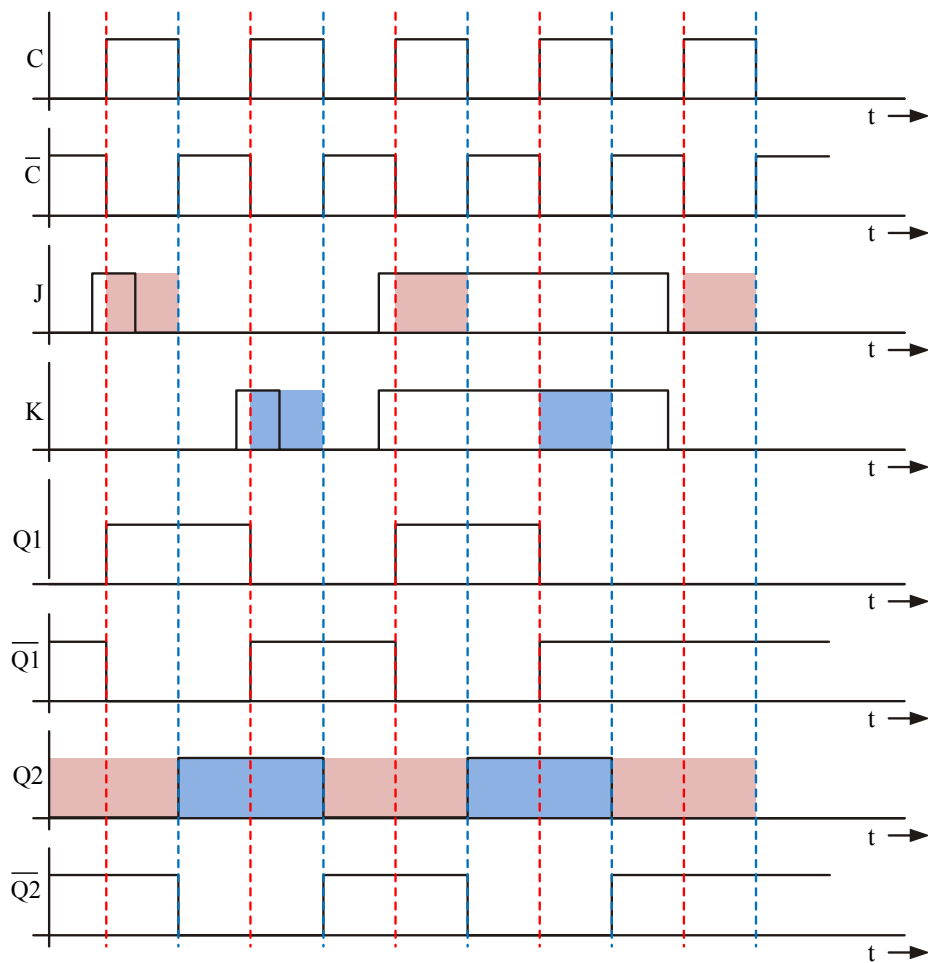


Bild 9.11: Signalverläufe eines taktzustandsgesteuerten JK-Master-Slave-Flip-Flop nach Bild 9.10.

Zeitspannen, in denen das Flip-Flop entweder auf eine Aktion am J- ($Q2 = 0$) oder K-Eingang ($Q2 = 1$) reagiert.

9.1.3 Taktflankengesteuerte Flip-Flop

Einflankengesteuerte D-Flip-Flop

Das einfachste einflankengesteuerte Flip-Flop ist ein D-Flip-Flop. Bild 9.12 zeigt das logische Schaltbild, den logischen Aufbau und ein Zeitdiagramm für ein solches D-Flip-Flop. Das Gesamt-Flip-Flop ist aus zwei taktzustandsgesteuerten D-Flip-Flops aufgebaut. Die am Eingang D anliegende Information wird während des L-Zustands des Taktsignals vom Master-Flip-Flop übernommen ($Q1$). Das Slave-Flip-Flop ist während dieser Zeit inaktiv. Beim Übergang des Taktes von L nach H wird das Slave-Flip-Flop aktiv und gleichzeitig das Master-Flip-Flop inaktiv. Dadurch wird die in diesem Moment an D anliegende logische Information im Master-Flip-Flop gespeichert und vom Slave-Flip-Flop an dessen Ausgang $Q2$ übergeben. Spätere Änderungen des Eingangssignals haben keinen Einfluss mehr auf den Ausgang, da der "Master" während des H-Pegels des Taktsignals gesperrt bleibt. Standardisierte D-Flip-Flops (z.B. 74ACT74) besitzen neben der Funktion des flankengesteuerten D-Flip-Flops zusätzlich die Funktion eines RS-Flip-Flops. Bild 9.13a zeigt das logische Symbol eines einflankengesteuerten D-Flip-Flops mit asynchronen Setz- und Rücksetzeingängen und Bild 9.13b die schaltungstechnische Realisierung mit Gattern. In der CMOS-Technik werden die Schaltungen auf eine möglichst geringe Anzahl von Transistoren optimiert. Eine tatsächliche Realisierung auf dem Chip kann deshalb völlig anders aussehen, als ein Grundentwurf auf Gatterebene. In Bild 9.13c ist eine Realisierung des beschriebenen einflankengesteuerten D-Flip-Flops mit taktunabhängigen Setz- und Rücksetzeingängen gezeigt. Als Anwender dieser integrierten Flip-Flop Bausteine muss man aber auch hier wieder beachten, dass für die taktunabhängigen Setz- und Rücksetzeingänge die gleiche Wahrheitstabelle wie für das RS-Flip-Flop mit NAND-Gattern nach Bild 9.4 gilt. Also gibt es auch in diesem Fall einen unerlaubten Eingangszustand: $\overline{R} = \overline{S} = 0$. Dieser Zustand ist in Bild 9.14 schraffiert gekennzeichnet.

Ein- und Zweiflankengesteuerte JK-Flip-Flops

Neben den taktzustandsgesteuerten Master-Slave Flip-Flops gibt es noch ein- und zweiflankengesteuerte Master-Slave Flip-Flops. Das unterschiedliche Verhalten eines taktzustandsgesteuerten und eines zweiflankengesteuerten JK-Flip-Flops ist in den Wahrheitstabellen und Zeit-Diagrammen in Bild 9.15 dargestellt. Die Übernahmzeit t_u , während der ein H-Pegel an den Eingängen $1J$ oder $1K$ im Master-Flip-Flop zwischengespeichert wird, ist beim zustandsgesteuerten Flip-Flop die gesamte Zeit, in der das Taktsignal $C1 = H$ ist (Bild 9.15a). Die Übernahmzeit t_u , während der ein H-Pegel an den Eingängen J oder K im Master-Flip-Flop zwischengespeichert wird, ist beim taktflankengesteuerten Flip-Flop nur die Zeit des Anstiegs der Flanke des Taktsignals (Bild 9.15b). Die Übergabe des gespeicherten Signals an den Ausgang erfolgt in

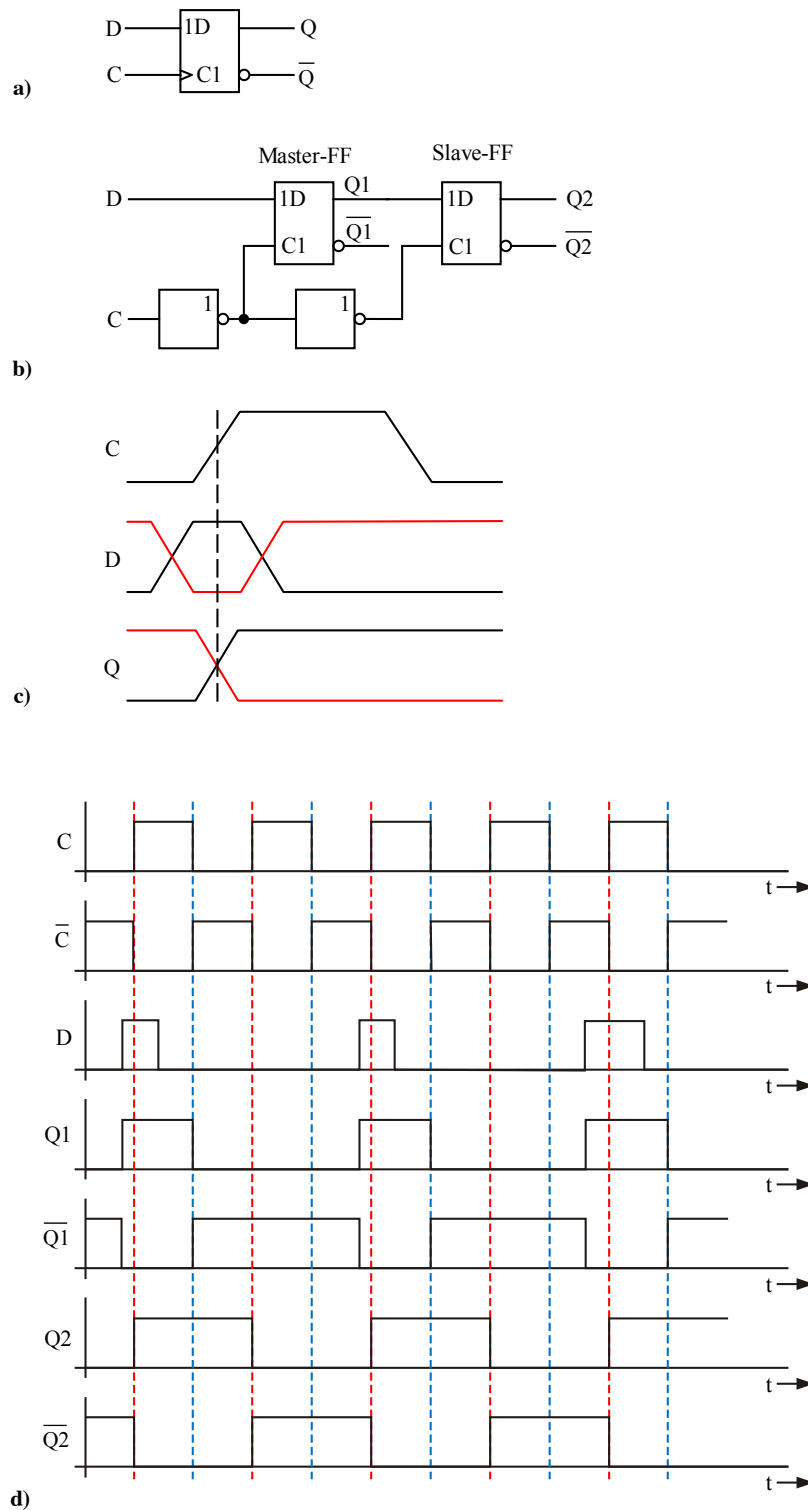


Bild 9.12: Einflankengesteuertes D-Flip-Flop a) Symbol, b) Schaltbild mit zwei zustandsgesteuerten D-Flip-Flops, c) Zeitabhängigkeit von Ein- und Ausgängen und d) Signalverläufe.

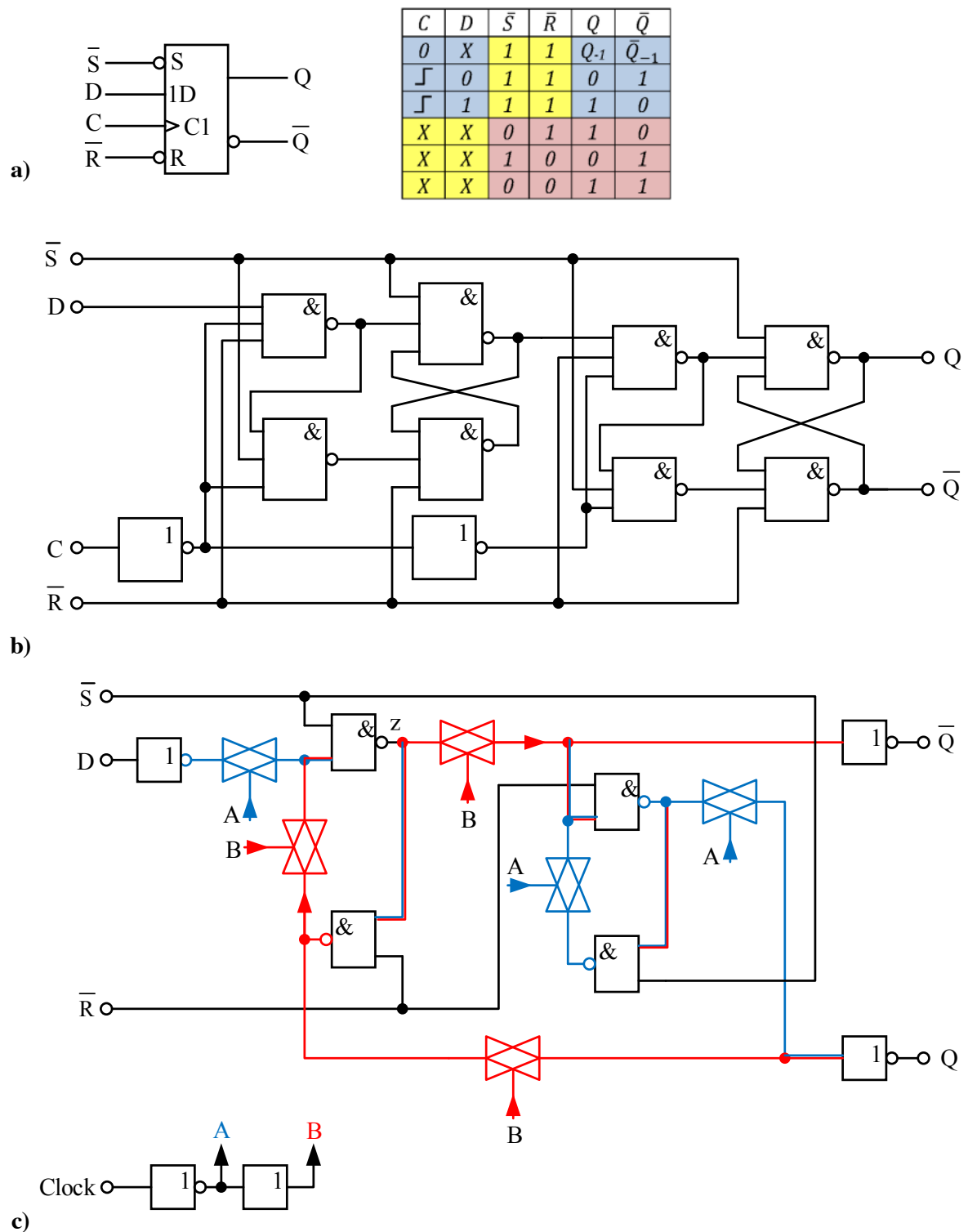


Bild 9.13: Einflankengesteuertes D-Flip-Flop mit Setz- und Rücksetzeingängen: a) Symbol und Wahrheitstabelle b) Schaltbild mit logischen Gattern und c) Realisierung der Schaltung in der CMOS-Technik

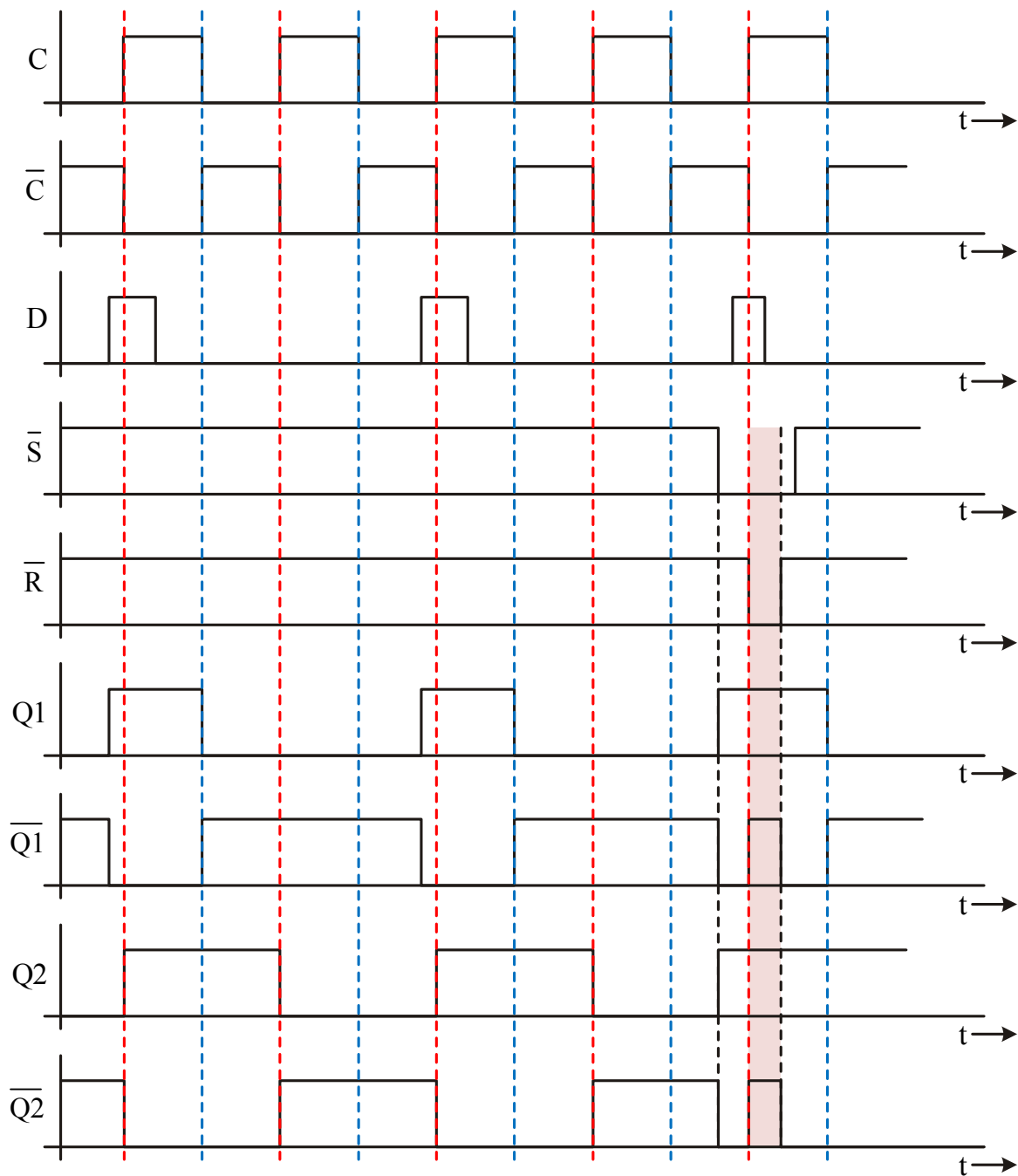


Bild 9.14: Signalverläufe eines einflankengesteuerten D-Flip-Flop mit Setz- und Rücksetzeingängen.

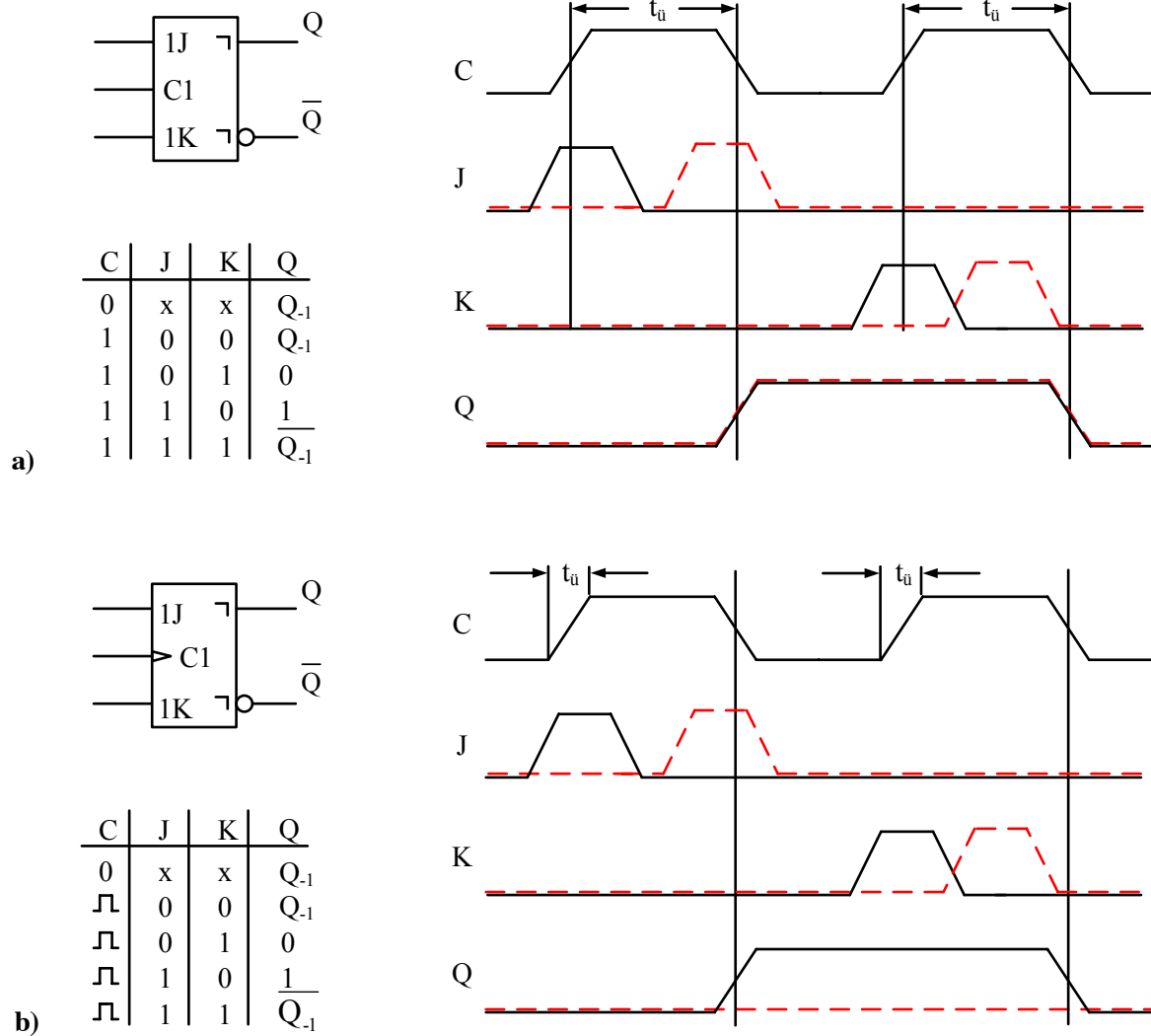


Bild 9.15: Logisches Symbol, zeitliche Verläufe der Ein- und Ausgangssignale und Wahrheitstabelle von JK-Flip-Flops: a) Zustandsgesteuertes JK-Flip-Flop, b) zweiflankengesteuertes JK-Flip-Flop.

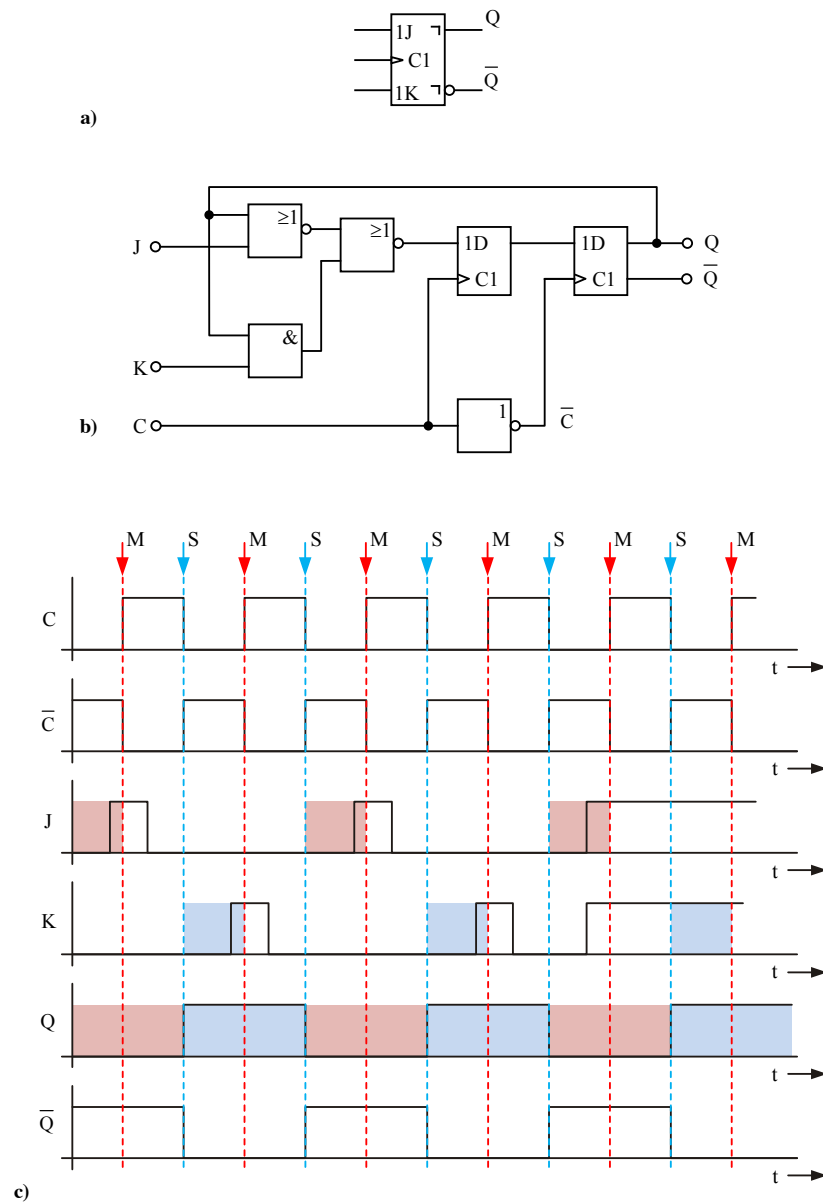


Bild 9.16: Zweiflankengesteuertes JK-Flip-Flop: a) logisches Symbol, b) Schaltbild mit zwei einflankengesteuerten D-Flip-Flops und Gattern und d) den Signalverläufen.

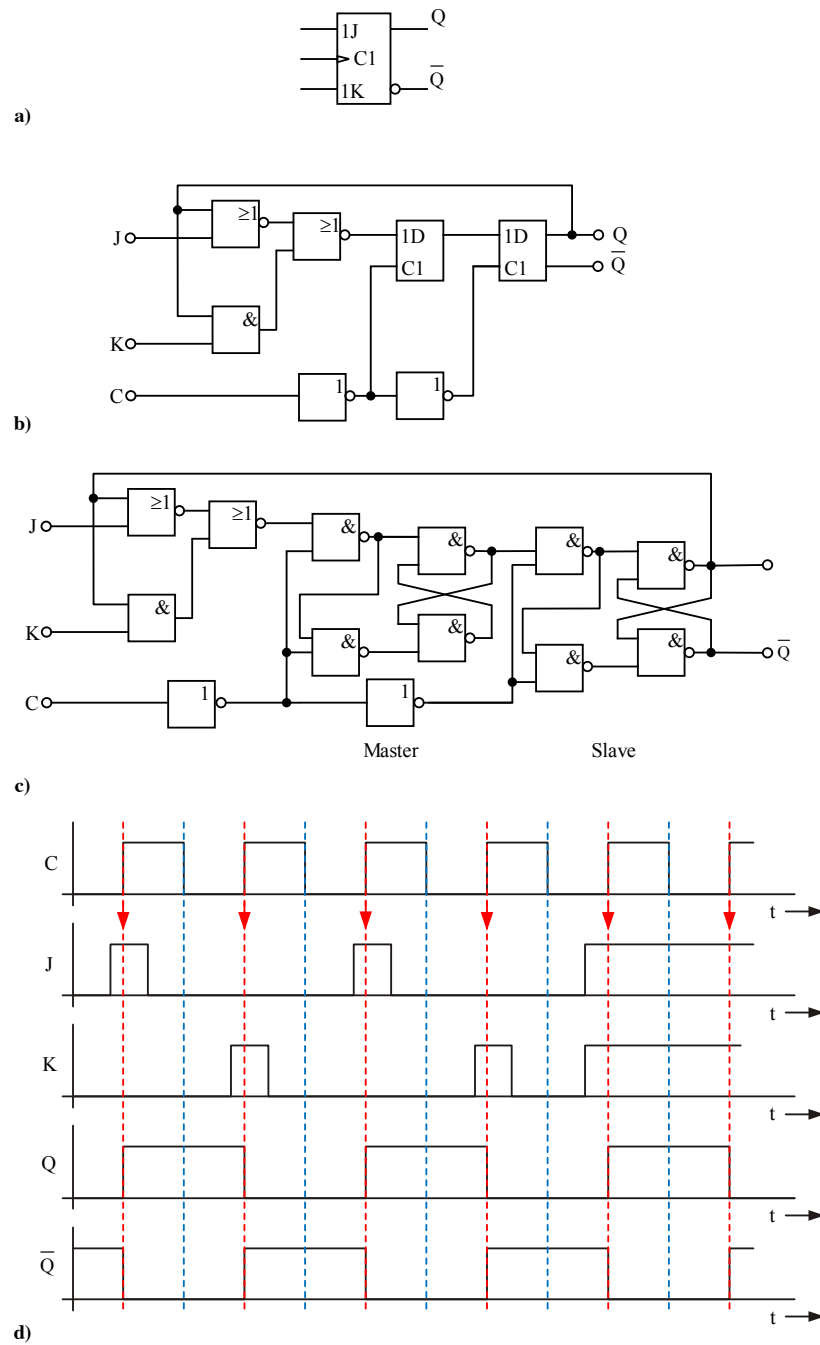


Bild 9.17: Einflankengesteuertes JK-Flip-Flop: a) logisches Symbol, b) Schaltbild mit zwei taktzustandsgesteuerten D-Flip-Flops und Gattern, c) Schaltbild mit logischen Gattern und d) die Signalverläufe.

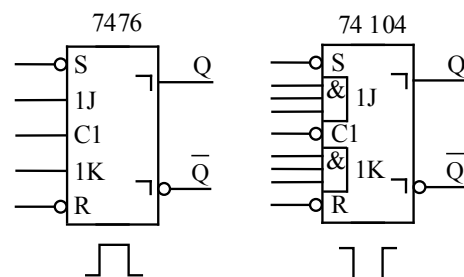
beiden Fällen erst beim Übergang des Taktsignals in den L-Pegel.

Bild 9.16 zeigt den Aufbau eines solchen zweiflankengesteuerten JK-Flip-Flops. Sowohl das Master- wie auch das Slave-Flip-Flop sind hier einflankengesteuerte D-Flip-Flops nach Bild 9.12. Bei der ansteigenden Flanke des Taktsignals wird das Eingangssignal in das Master-Flip-Flop übernommen. Bei der abfallenden Flanke des Taktsignals erhält durch den vorgeschalteten Inverter das Slave-Flip-Flop eine ansteigende Taktflanke und übernimmt die Ausgangsinformation des Master-Flip-Flops.

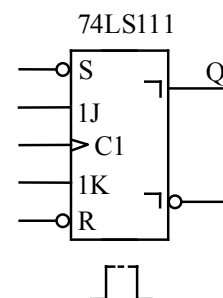
Bei der heute überwiegend eingesetzten Ausführungsform des JK-MasterSlave Flip-Flops erfolgt die Übernahme des Eingangssignals ebenfalls während der Flanke des Taktsignals, d.h. beim Wechsel von einem logischen Zustand in den anderen. Die Übergabe der so gespeicherten Information erfolgt aber nicht um die Dauer des Taktsignals verzögert, sondern sofort. Das Schaltungssymbol eines einflankengesteuerten JK-Flip-Flops und der Aufbau mit Gattern bzw. mit D-Flip-Flops und Gattern ist in Bild 9.17 zu sehen. Der Kern des Flipflops wird hier von zwei taktzustandsgesteuerten bzw. einem einflankengesteuerten D-Flipflop gebildet.

Die nachfolgende Zusammenstellung zeigt logische Symbole verschiedener Ausführungen von JK-Flip-Flops der Baureihe 74... und beschreibt in Stichworten deren Funktion. Zusätzlich zu den bisher beschriebenen Funktionen der JK-Flip-Flops besitzen alle gezeigten Bausteine noch asynchrone Setz- und Rücksetzeingänge (S und R). Wird z.B. an den SET-Eingang zu einem beliebigen Zeitpunkt ein logischer LOW-Pegel angelegt, so wird unmittelbar danach am Ausgang Q ein HIGH-Pegel und an $\bar{Q}$ ein LOW-Pegel anliegen. Zu beachten ist jedoch, dass zu keinem Zeitpunkt an beiden Eingängen gleichzeitig ein LOW-Pegel angelegt werden darf.

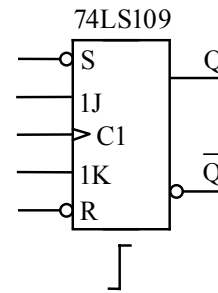
Übernahme einer Information während der Dauer des H bzw. L-Pegels des Taktes. Verzögerte Abgabe der gespeicherten Eingangsinformation beim Übergang des Taktsignals nach L bzw. H.



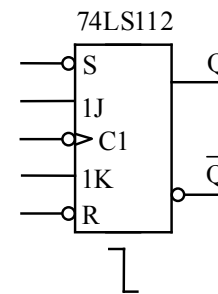
Übernahme nur während der ansteigenden Flanke des Taktes. Verzögerte Abgabe bei der fallenden Flanke des Taktes.



Übernahme nur während der ansteigenden Flanke des Taktes. Unverzögerte Abgabe an den Ausgang.



Übernahme nur während der fallenden Flanke des Taktes. Unverzögerte Abgabe an den Ausgang.



9.2 Zähler

Lernziele

- Kennenlernen und Verstehen der Funktion und des Aufbaus von Zählerschaltungen
 - asynchronen Zählerschaltungen
 - synchronen Zählerschaltungen
 - binäre und BCD-Zähler
- Kennenlernen der Einsatzbereiche von Zählern und Frequenzteilern

Zähler und Frequenzteiler sind Bestandteile von nahezu allen digitalen Systemen. Zähler werden nach der Systemarbeitsweise, also synchrone bzw. asynchrone Taktung, dem Zahlenkode und ihrer Zählrichtung unterschieden. Man spricht vom asynchronen Betrieb, wenn die Flipflops zu unterschiedlichen Zeitpunkten geschaltet werden. Wenn alle Flip-Flop zum gleichen Zeitpunkt geschaltet werden, sprechen wir von synchronen Zählern. Nachfolgend werden die wichtigsten Zählerarten, aufgebaut mit den Flip-Flop-Schaltungen des vorangegangenen Kapitels, vorgestellt und detailliert beschrieben.

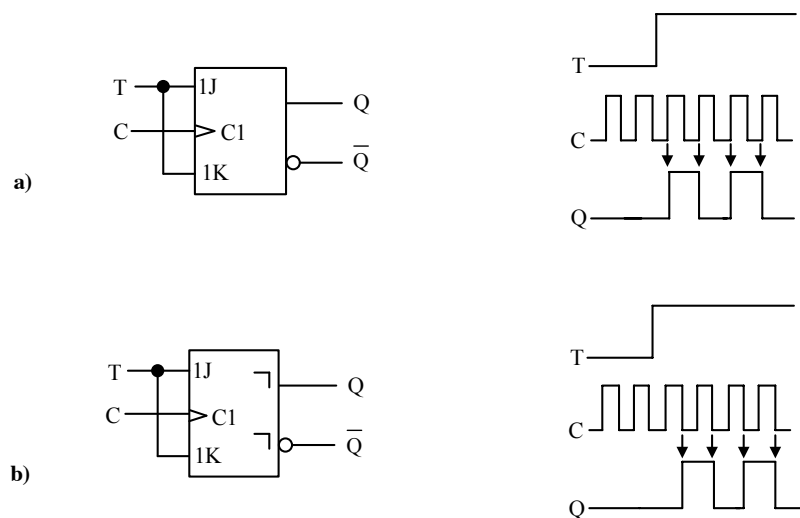


Bild 9.18: Aufbau und Funktion eines T-Flip-Flop: a) einflankengesteuertes JK-Flip-Flop, b) zweiflankengesteuertes JK-Flip-Flop.

9.2.1 Asynchrone Zähler

9.2.2 Asynchrone Dualzähler

Dual-Vorwärtszähler

Für Zähler werden häufig einflankengesteuerte T-Flip-Flops eingesetzt, die einfach aus JK-Flip-Flops gebaut werden können. In Bild 9.18 ist gezeigt, wie man ein T-FF aus einem JK-Flip-Flop ableiten kann. Durch die feste Verbindung des J- und des K-Eingangs kann nur $J = K = 0$ oder $J = K = 1$ anliegen. Aus der Wahrheitstabelle des JK-Flip-Flops ist zu entnehmen, dass im ersten Fall die Ausgänge unverändert bleiben; das Flip-Flop befindet sich im inaktiven Zustand. Im zweiten Fall wird der Ausgangspegel bei jedem neuen Eintreffen einer Flanke am Takteingang invertiert, d.h. war $Q = 0$, wird $Q = 1$ und umgekehrt. Bild 9.18a zeigt die Schaltung und das zeitliche Verhalten eines einflankengesteuerten T-Flip-Flops und Bild 9.18b das, eines zweiflankengesteuerten T-Flip-Flops.

Bild 9.19a zeigt die Schaltung eines 4-Bit-Dual-Vorwärtszählers. Die von den Ausgängen gebildete Dualzahl ergibt sich aus der Wichtung der Ausgänge: $Q_A = 2^0$; $Q_B = 2^1$; $Q_C = 2^2$; $Q_D = 2^3$. Damit kann der 4-Bit-Zähler von 0 bis 15 zählen. Nach Erreichen der maximal möglichen Zahl schalten alle Flip-Flops wieder auf null zurück und der Zählvorgang beginnt von neuem. Die Flip-Flops im Bild 9.19 schalten mit der ansteigenden (positiven) Signalfanke. Deshalb müssen die einzelnen Stufen mit dem invertierten Ausgang Q angesteuert werden. Das zugehörige Zeitablaufdiagramm (Bild 9.19b) macht dies ersichtlich. Die Gatterlaufzeiten sind im Diagramm nicht berücksichtigt.

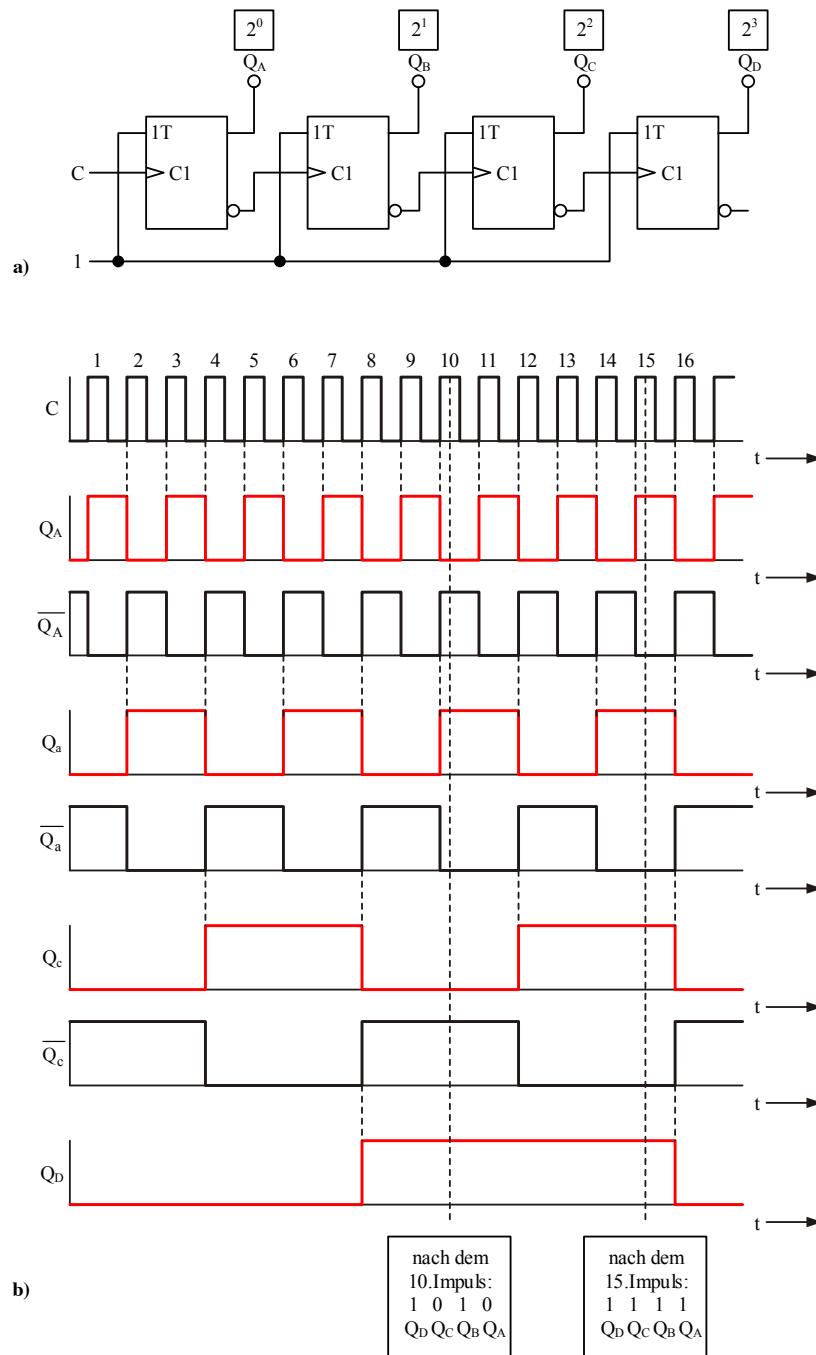


Bild 9.19: Asynchroner Vorwärtszähler mit einflankengesteuerten T-Flip-Flops.

Die Realisierung eines dualen 4-Bit-Vorwärtszählers mit zweiflankengesteuerten T-Flip-Flop, d.h. das Master-FF übernimmt die Information mit der ansteigenden Signalflanke und das Slave-FF schaltet mit der abfallenden Flanke, mit dem zugehörigen Zeitblaufdiagramm zeigt Bild 9.20. Durch die Beschaltung der T-Eingänge mit 1 (High-Pegel) arbeiten die Flip-Flops immer im Toggle-Modus. Im Gegensatz zu den Vorwärts-

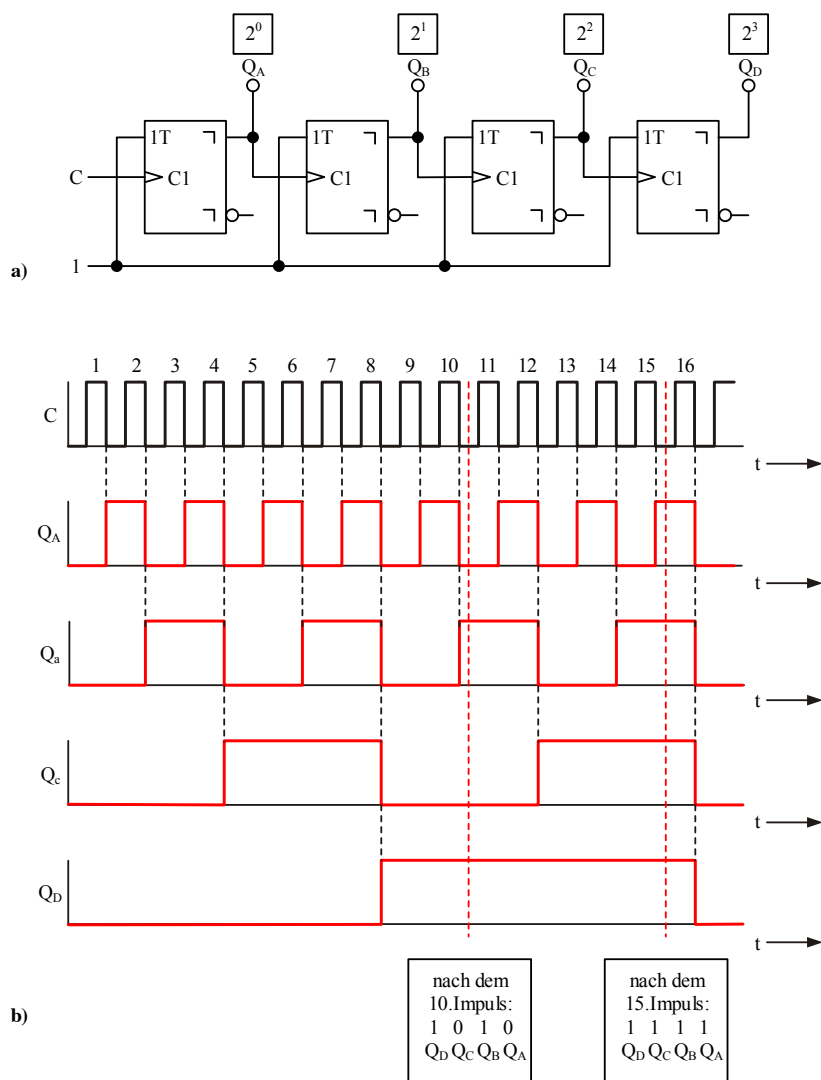


Bild 9.20: Asynchroner Vorwärtszähler mit zweiflankengesteuerten T-Flip-Flops.

zählern zählen duale Rückwärtszähler von ihrem möglichen Höchstwert ab rückwärts bis auf null und beginnen dann wieder bei ihrem Höchstwert erneut rückwärts zu zählen. In Bild 9.21 ist die Schaltung und das Zeitdiagramm eines 4-Bit-Dual-Rückwärtszählers mit einflankengesteuerten T-Flip-Flops gezeigt. Man kann sehen, dass er sehr einfach aus dem Vorwärtszähler gebaut werden kann. Man muss nur die nichtinvertierten Ausgänge Q direkt mit dem Takteingang des folgenden T-Flip-Flops verbinden. Dadurch erreicht man, dass nach der ersten Impulsflanke alle Ausgänge auf 1 gesetzt werden. Damit beginnt der Zähler von seinem maximalen binären Wert 1111 (dezimal 15) mit dem Zählvorgang. Mit der positiven Flanke des 16. Impulses hat der Zähler den binären Wert 0000 erreicht. Mit dem nächsten Impuls beginnt der Zählvorgang erneut bei dem Wert 1111.

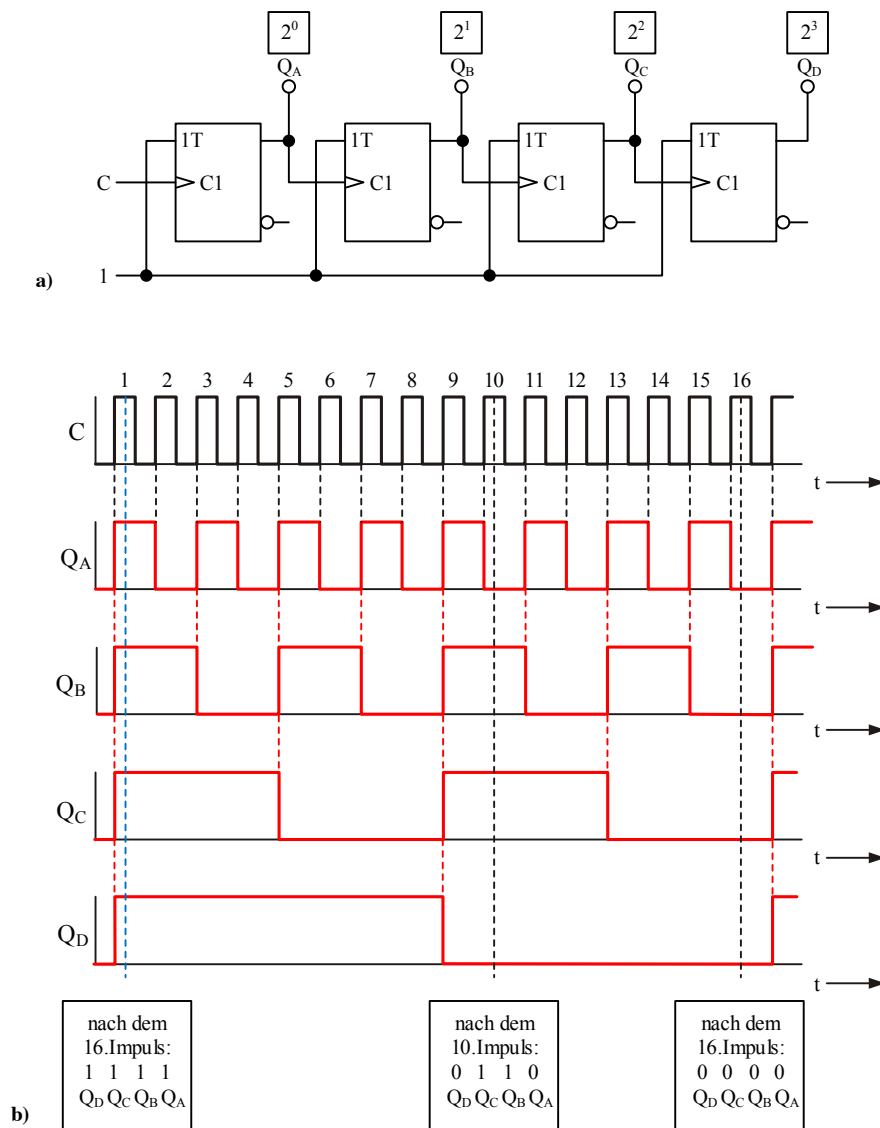


Bild 9.21: Asynchroner Rückwärtszähler mit einflankengesteuerten T-Flip-Flops.

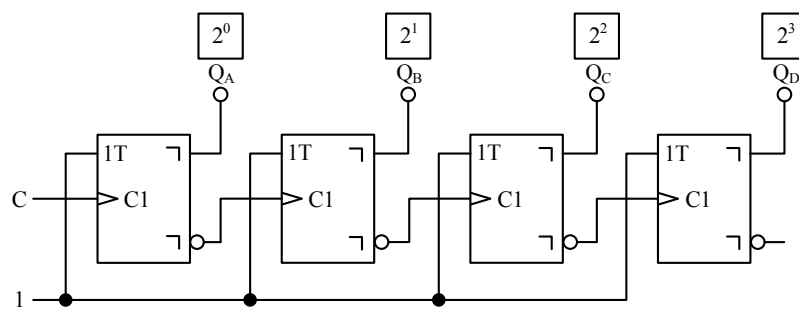


Bild 9.22: Asynchroner Rückwärtszähler mit zweiflankengesteuerten T-Flip-Flops.

Das gleiche Zeitdiagramm erhält man auch, wenn man zweiflankengesteuerte T-Flip-Flops einsetzt (Bild 9.22) und anstelle der Q -Ausgänge die $\overline{Q}$ -Ausgänge zur Ansteuerung der nachfolgenden Stufen verwendet. Der wesentliche Vorteil von asynchronen Dualzählern ist der einfache Schaltungsaufbau mit einfachen T-Flip-Flops. Der Nachteil dieses Schaltungsaufbaus soll nun etwas genauer betrachtet werden: Bei den bisher dargestellten Zeitdiagrammen der asynchronen Dualzähler haben wir die Gatterlaufzeit der T-Flip-Flops vernachlässigt, da für uns bisher das Ergebnis des Zählens entscheidend war. Deshalb soll anhand eines Beispiels der Einfluss der Gatterlaufzeit auf das Zählergebn näher betrachtet werden. Dazu nehmen wir einen Ausschnitt aus dem Zeitdiagramm des Rückwärtszählers nach Bild 9.21b) und betrachten den Übergang von 0000 nach 1111, wie in Bild 9.23 dargestellt. Zwischen dem Takteingang und dem Q -Ausgang jedes einzelnen Flip-Flops entsteht durch die Signallaufzeiten in den einzelnen Gattern des Flip-Flops eine Verzögerungszeit t_{pd} . Am Eingang des ersten Flip-Flops liegt der Takt an. Das Signal am Ausgang erscheint nun um $1 \cdot t_{pd}$ gegenüber dem Taktsignal verzögert. Da der Ausgang des ersten Flip-Flops den Takteingang des zweiten Flip-Flops ansteuert, erscheint dessen Ausgangssignal wiederum um $1 \cdot t_{pd}$ später, also gegenüber dem Eingangstaktsignal um $2 \cdot t_{pd}$ verzögert. Der Ausgang der dritten Stufe des Zählers ist dann bereits um $3 \cdot t_{pd}$ gegenüber dem Eingangstaktsignal verzögert. Da sich dies bei jeder weiteren Stufe genauso fortsetzt, wird die Zeitspanne, während der das Ergebnis des Zählvorgangs exakt abgelesen und weiterverarbeitet werden kann, immer kleiner. Erhöht man die Taktfrequenz des Eingangstaktes, reduziert sich ebenfalls die Ablesezeitspanne. Deshalb muss beim Einsatz von asynchronen Dualzählern dieses Verhalten vom Anwender berücksichtigt werden.

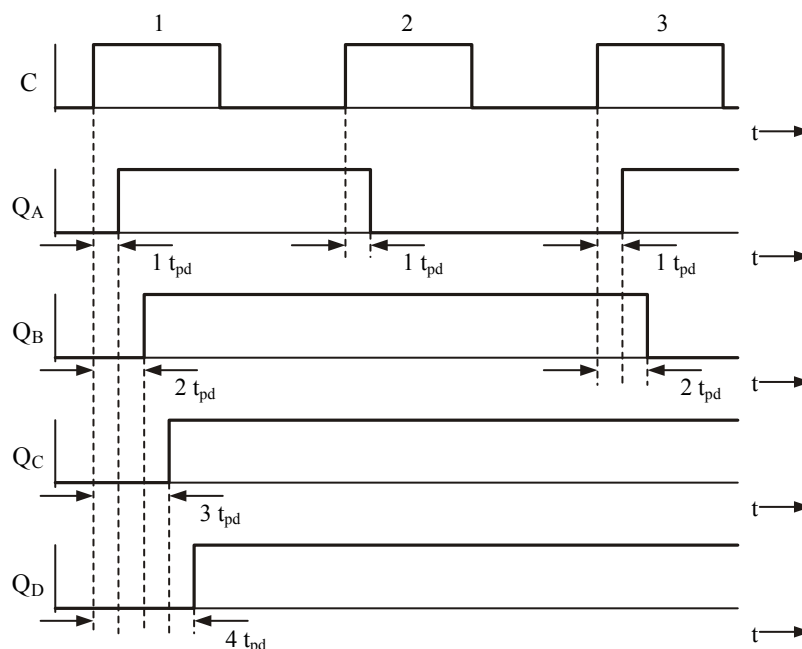


Bild 9.23: Einfluss der Gatterlaufzeiten der T-Flip-Flops beim asynchronen Dualzähler aus einflankengesteuerten JK-Flip-Flops.

BCD-Zähler

Bisher wurden die verschiedenen Schaltungsvarianten asynchroner dualer Zähler abgehandelt. Die Wertigkeit der einzelnen Ausgänge ist dabei immer eine Potenz der Zahl 2. Will man jedoch auf das dezimale Zahlensystem übergehen, so müssen die Schaltungen entsprechend verändert werden, da diese Zähler nach dem 10. Taktimpuls nicht mehr weiterzählen dürfen, sondern bei 0 wieder neu mit dem Zählen beginnen müssen. Zur Veranschaulichung der Funktionsweise sollen zunächst die Zustandstabelle (Tabelle 9.1) und das Zeitdiagramm in Bild 9.24 betrachtet werden. Auch hier wird die Gatterlaufzeit wieder vernachlässigt. Im Ausgangszustand soll der Zähler auf 0 stehen. Eine Rechteckschwingung am Takteingang liefert die Zählimpulse. Wie schon bei Dualzählern wird angenommen, dass der Zähler aus zweiflankengesteuerten JK-Flip-Flops aufgebaut ist, die bei der ansteigenden Flanke des Taktes die Eingangsinformation in das Master-Flip-Flop übernehmen und diese bei der abfallenden Flanke über das Slave-Flip-Flop an den Ausgang weiterleiten. Aus der Zustandstabelle und den Zeitdiagramm erkennt man, dass sich die Ausgänge $Q_D \dots Q_A$ während der ersten 9 Taktimpulse wie beim Dualzähler verhalten. Beim zehnten Impuls jedoch tritt an den Ausgängen $Q_D \dots Q_A$ nicht der binäre Zustand 1010 auf, sondern 0000; d.h. der Zähler wird zurückgesetzt und es wird neu mit dem Zählen begonnen. Dies erfordert gegenüber dem Dualzähler eine Erweiterung der Schaltung. Zunächst könnte man sich vorstellen, JK-Flip-Flops mit einem Rücksetzeingang einzusetzen. Durch eine Überprüfung des Ausgangszustands $10 = 1\ 0\ 1\ 0$ und ein sofortiges Löschen der Ausgänge Q_D und Q_B über den Reset-Eingang würde auch wieder der Wert $0\ 0\ 0\ 0$ an den Ausgängen anliegen jedoch würde kurzzeitig auch der Wert $1\ 0\ 1\ 0$ abzulesen sein. Abhängig vom Zeitpunkt des Ablesens des Zählerstandes ist damit kurzzeitig ein Ablesefehler möglich. Um dies zu verhindern, muss eine andere Lösung gefunden werden. Die Zustandstabelle und das Zeitdiagramm zeigen, dass zumindest bis zur Zahl 8 kein Unterschied zwischen Dual- und BCD-Zähler besteht. Also können die ersten drei JK-Flip-Flops wie beim Dualzähler verschaltet werden. Nach den siebten Taktimpuls sind die Ausgänge Q_A , Q_B und $Q_C = 1$. Aus dem Zeitdiagramm kann man ablesen, dass das vierte Flip-Flop nur gesetzt werden darf, wenn wenigstens Q_B und $Q_C = 1$ sind. Dieser Zustand kann mit dem Q -Ausgang des ersten Flip-Flops zum Zeitpunkt t_1 in das vierte Flip-Flop übernommen werden, wenn $J = 1$ ist. Dies erreicht man durch eine AND-Verknüpfung von Q_B und Q_C und eine Verbindung von Q_A mit dem Takteingang dieses Flip-Flops. Zu prüfen ist jetzt nur noch, ob das vierte Flip-Flop nach dem zehnten Taktimpuls auch auf 0 zurückgesetzt wird. Bei der Übernahme von J zum Zeitpunkt t_2 ist Q_B und $Q_C = J = 0$. Der K-Eingang liegt fest auf logisch 1. Mit der ansteigenden Flanke von Q_A am Takteingang des letzten Flip-Flops wird damit das Master-Flip-Flop auf $Q = 0$ gesetzt. Dieser Wert wird bei der fallenden Flanke an den Ausgang Q_D übergeben. Damit sind alle vier Ausgänge des Zählers $Q_A \dots Q_D = 0$. Da Q_A aber auch das zweite Flip-Flop ansteuert, muss verhindert werden, dass dieses gesetzt werden kann. Durch eine Verbindung des J-Eingangs dieses Flip-Flops mit dem Q -Ausgang des letzten Flip-Flops kann dies gewährleistet werden, denn sobald $Q_D = 1$ wird, ist $J = 0$ und damit ist ein Setzen des Ausgangs Q_B nicht mehr möglich. Bild 9.25 zeigt die nach obiger Beschreibung aufgebaute Schaltung. Der RCO-Ausgang (Ripple

Takt (C)	Q_D	Q_C	Q_B	Q_A	Dezimalwert N
0,1	Q_{D-1}	Q_{C-1}	Q_{B-1}	Q_{A-1}	N_{-1}
	0	0	0	0	0
	0	0	0	1	1
	0	0	1	0	2
	0	0	1	1	3
	0	1	0	0	4
	0	1	0	1	5
	0	1	1	0	6
	0	1	1	1	7
	1	0	0	0	8
	1	0	0	1	9
	0	0	0	0	10
	0	0	0	1	11

Tabelle 9.1: Zustandstabelle eines BCD-Zählers.

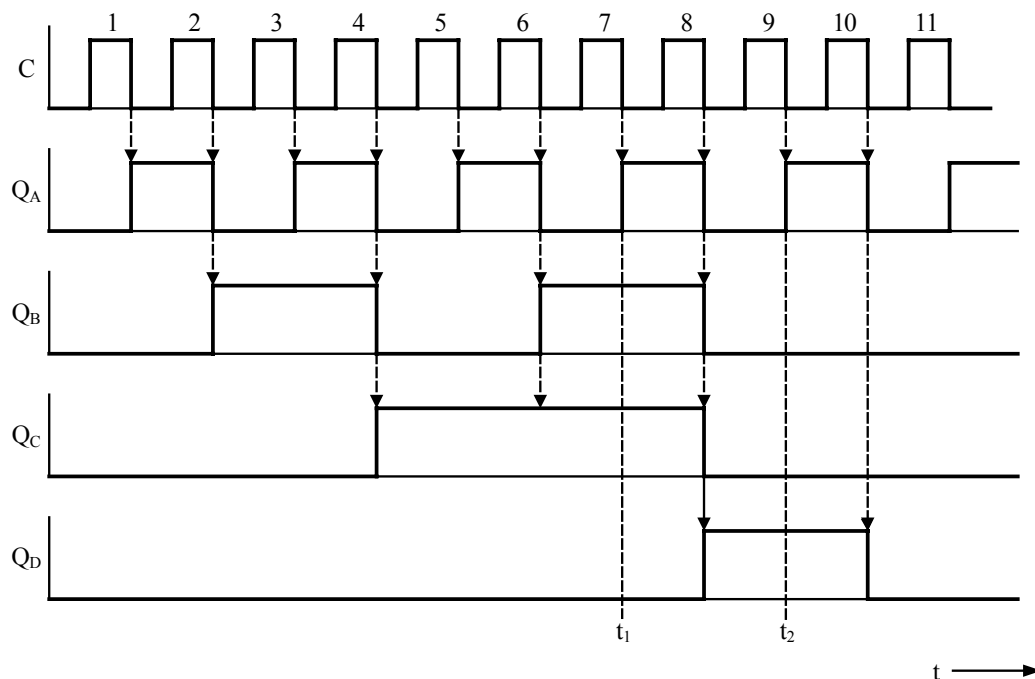


Bild 9.24: Zeitdiagramm eines BCD-Zählers mit zweiflankengesteuerten JK-Flip-Flops.

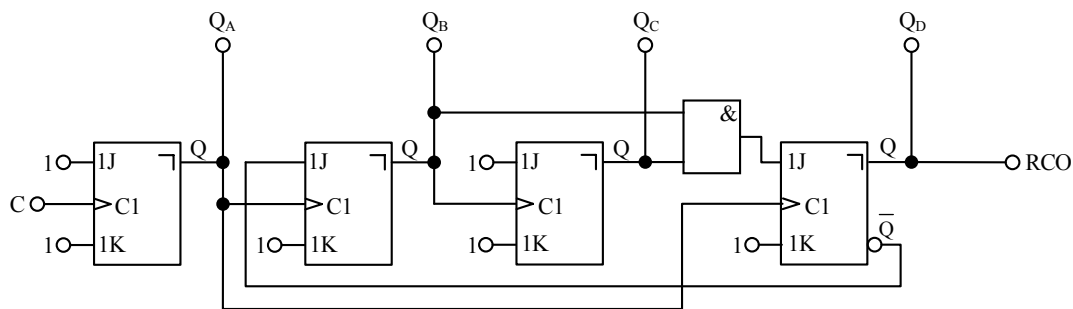


Bild 9.25: Schaltung eines asynchronen BCD-Zählers für eine Dezimalstelle mit zweiflankengesteuerten JK-Flip-Flops.

(Carry Out) entspricht beim asynchronen BCD-Zähler dem Ausgang Q_D und dient zur Kaskadierung solcher BCD-Zähler zu mehrstelligen Dezimalzählern.

9.2.2.1 Asynchrone Modulo- n -Zähler

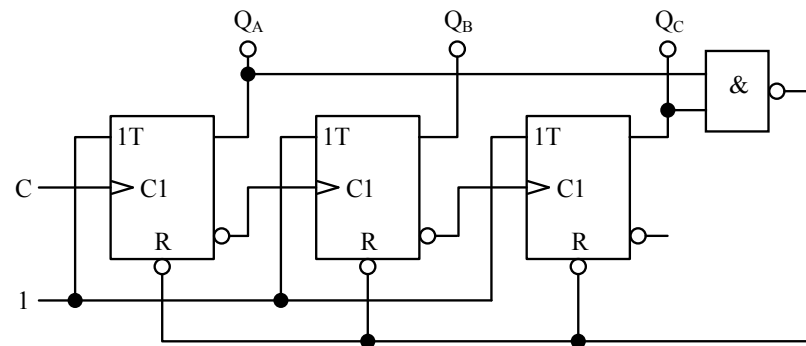
Für unterschiedliche Aufgaben in der Störungstechnik und z.B. bei der Zeitmessung werden Zähler benötigt, die bis zu einem gewünschten Zahlenwert zählen, dann auf 0 zurücksetzen und die Zählung erneut beginnen oder aber bei diesem Wert stehen bleiben und auf ein neues Startsignal warten. Die Zahl, bis zu der zu zählen ist, sollte dabei beliebig sein. Diese Zähler werden Modulo- n -Zähler genannt (von Modulus - lateinisch - Maß). Der kleine Buchstabe n steht für die Anzahl der möglichen Zählerzustände. So ist z.B. ein BCD-Zähler ein Modulo-10-Zähler. Er zählt zwar nur bis 9, aber einschließlich der Zahl 0 hat er insgesamt 10 mögliche Zählerzustände.

Modulo-5-Zähler

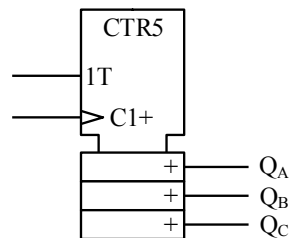
Ein Modulo-5-Zähler muss bis 4 zählen können, und mit dem 5. Impuls auf 0 gesetzt werden. Die Schaltung des Zählers ist in Bild 9.26a gezeigt. Der Zähler wird wie der asynchrone Dualzähler aus T-Flip-Flops aufgebaut. Für die Realisierung eines Modulo-5-Zählers werden 3 Flip-Flops benötigt. Dieser Zähler kann von 0 bis 7 zählen. Beim Übergang von 4 auf 5 muss der Zähler auf 0 gestellt werden. Das Rückstellen kann auf die gleiche Art wie beim BCD-Zähler erreicht werden. Wenn $Q_A = 1$ und $Q_C = 1$ sind, soll der Zähler zurückgestellt werden. Als Rückstellsignal wird ein 0-Signal benötigt. Die Ausgänge Q_A und Q_C werden über ein NAND-Gatter verknüpft. Der Ausgang des NAND-Gatters liefert das Rückstellsignal 0, wenn Q_A und Q_C logisch High-Pegel führen. Das Schaltsymbol eines Modulo-5-Zählers ist in Bild 9.26b angegeben.

Modulo-60-Zähler

Ein Modulo-60-Zähler wird z.B. für elektronische Uhren benötigt. Die Sekunden werden von 0 bis 60 gezählt. Wie viele Flip-Flops werden dafür benötigt? Mit 5 Flip-Flops



a)



b)

Bild 9.26: a) Schaltung und b) Schaltsymbol eines Modulo-5-Zählers.

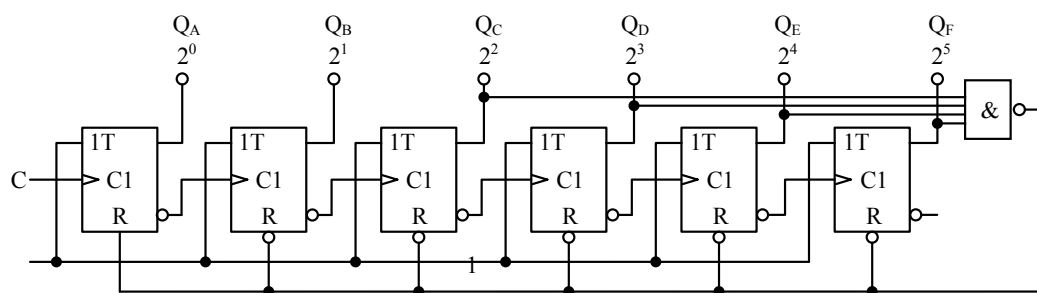


Bild 9.27: Schaltung eines asynchronen Modulo-60-Zählers.

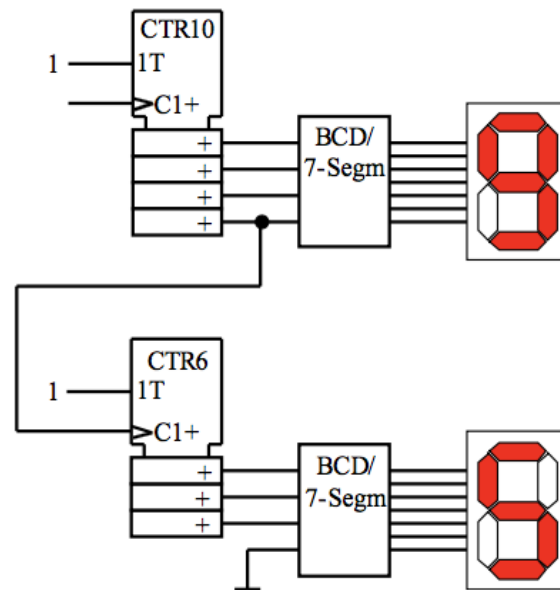


Bild 9.28: Schaltung aus Modulo-10-Zähler und Modulo-6-Zähler mit angeschlossener 7-Segment-Anzeige.

kann man bis 31 zählen, mit 6 Flip-Flops bis 63. Damit benötigen wir also 6 Flip-Flops. Bild 9.27 zeigt das Schaltbild eines Modulo-60-Zählers mit T-Flip-Flops. Die T-Flip-Flops werden wiederum aus JK-Flip-Flops gewonnen. Beim Erscheinen des Dezimalwertes 60 muss der Zähler auf 0 zurückgestellt werden. Dafür müssen die Ausgänge Q_C , Q_D , Q_E und Q_F High-Pegel führen. Aus diesen Signalen wird dann das Rückstell-signal gewonnen. Der Zähler ist für die Sekunden-zählung gut geeignet, wenn die Sekunden nicht als Dezimalzahl angezeigt werden sollen. Sollen die Sekunden als Dezimalzahl angezeigt werden, ist es zweckmäßig, Einer und Zehner getrennt zu zählen. Für die Einer benötigt man dann einen Modulo-10-Zähler (BCD-Zähler), für die Zehner einen Modulo-6-Zähler. Das Schaltbild dieses zusammengesetzten Zählers ist in Bild 9.28 dargestellt. Die Ausgangssignale dieser beiden Zähler können einem BCD-7-Segment Decoder zugeführt und als Dezimalziffern mit 7 Segment-Anzeigen dargestellt werden.

Modulo-13-Zähler mit Wartepflicht

Im Folgenden soll ein Modulo- n -Zähler aufgebaut und besprochen werden, der bei Erreichen des Dezimalwertes 12 stehen bleibt, wartet und an einem Ausgang Z ein High-Signal bereitstellt. Der Zähler soll auf Tastendruck zurückstellen und dann erneut mit dem Zählvorgang beginnen. Für eine solche Zählerschaltung werden 4 Flip-Flops benötigt. Der Eingang E muss über ein AND-Glied sperrbar sein. Als Sperrsignal wird Low-Pegel verwendet. Das Sperrsignal wird aus den Ausgangssignalen Q_C und Q_D mit Hilfe eines NAND-Gatters gewonnen. Die Schaltung des Zählers ist in Bild 9.29 gezeigt. Bei Erreichen des Dezimalwertes 12 ($Q_C = 1, Q_D = 1$) liegt am Ausgang X des NAND-Gatters logisch 0 an. Der Eingang sperrt; damit bleibt der Zähler stehen.

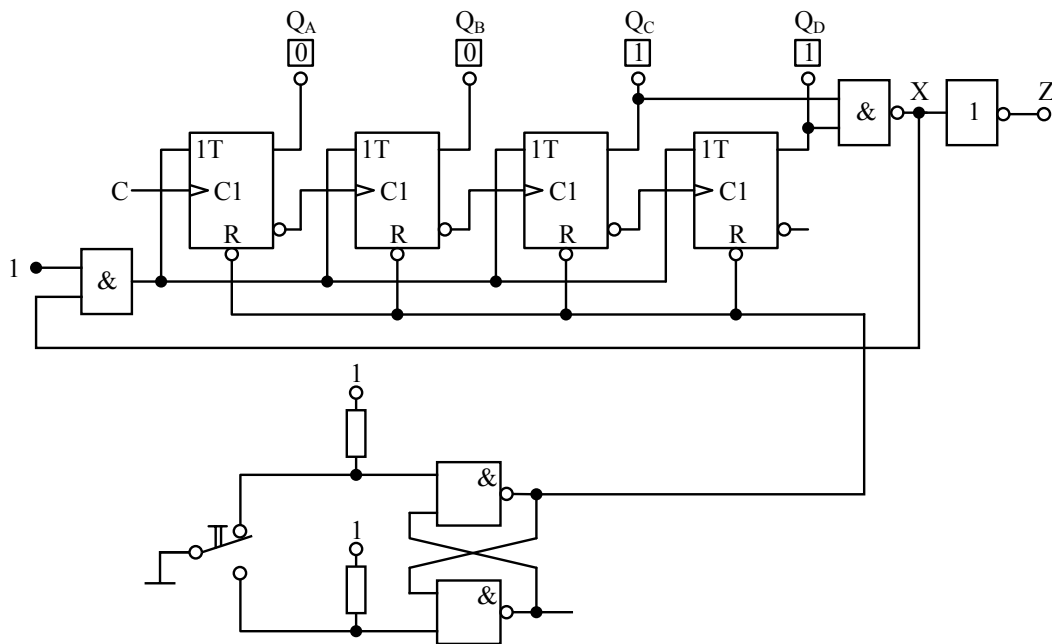


Bild 9.29: Modulo-13-Zähler mit Wartepflicht.

Gleichzeitig erscheint am Ausgang Z ein Hochpegel. Durch Drücken der Taste wird der Zähler zurückgestellt. Die Eingangssperre wird aufgehoben, da Q_C und Q_D jetzt 0-Pegel führen. Der Zähler beginnt mit einem neuen Zählvorgang. Da 13 Zählerzustände einschließlich null möglich sind, ist dieser Zähler ein Modulo-13-Zähler.

9.2.2.2 Synchrone Zähler

Synchrone Dualzähler

Bei den asynchronen Zählern wird, wie beschrieben, jede folgende Stufe vom Ausgang der vorhergehenden Stufe direkt angesteuert, mit dem Nachteil, dass sich die Gatterlaufzeiten der einzelnen Flip-Flops addieren. Um dies zu vermeiden, hat man synchrone Zähler entwickelt. Bei den synchronen Zählern liegt der Takt gemeinsam an allen Stufen an. Damit aber ein korrektes Zählen möglich ist, muss die Ansteuerung der Flip-Flops nun über die T-Eingänge erfolgen. Durch $T = 0$ oder $T = 1$ muss jetzt festgelegt werden, welche der Flip-Flops beim Eintreffen des Taktsignals bzw. der Taktflanke toggeln (umschalten) dürfen. Einen guten Überblick liefert auch hier die Zustandstabelle für einen Vorwärtszähler. Nach dem Ausgangszustand 0 0 0 0 darf nur das erste Flip-Flop toggeln, alle anderen nicht. Demzufolge darf auch nur der T-Eingang dieses Flip-Flops auf $T = 1$ sein, während die restlichen Flip-Flops mit $T = 0$ angesteuert werden müssen. Beim nächsten Takt müssen dann die Flip-Flops 1 und 2 umschalten, dann nur Flip-Flop 1, danach die Flip-Flops 1 bis 3 usw. Für die T-Eingänge der vier Flip-Flops ergeben sich damit die in den rechten Spalten der Zustandstabelle 9.2 angegebenen

Takt (C)	Q_D	Q_C	Q_B	Q_A	Dezimalwert N	T_D	T_C	T_B	T_A
0,1	Q_{D-1}	Q_{C-1}	Q_{B-1}	Q_{A-1}	N_{-1}	x	x	x	x
	0	0	0	0	0	0	0	0	1
	0	0	0	1	1	0	0	1	1
	0	0	1	0	2	0	0	0	1
	0	0	1	1	3	0	1	1	1
	0	1	0	0	4	0	0	0	1
	0	1	0	1	5	0	0	1	1
	0	1	1	0	6	0	0	0	1
	0	1	1	1	7	1	1	1	1
	1	0	0	0	8	0	0	0	1
	1	0	0	1	9	0	0	1	1
	1	0	1	0	10	0	0	0	1
	1	0	1	1	11	0	1	1	1
	1	1	0	0	12	0	0	0	1
	1	1	0	1	13	0	0	1	1
	1	1	1	0	14	0	0	0	1
	1	1	1	1	15	1	1	1	1

Tabelle 9.2: Zustandstabelle eines dualen Vorwärtzzählers.

logischen Zustände. Man erkennt sofort, dass der Eingang T_B des zweiten Flip-Flops exakt dem Zustand des Ausgangs Q_A des ersten Flip-Flops entspricht. Der Eingang T_C des dritten Flip-Flops entspricht einer UND-Verknüpfung der Ausgänge Q_A und Q_B und der Eingang T_D des vierten Flip-Flops einer UND-Verknüpfung der Ausgänge Q_A , Q_B und Q_C . Damit lässt sich durch Hinzufügen von UND-Gattern die Steuerung der einzelnen Flip-Flops vornehmen. Man spricht auch oft davon, dass die Flip-Flops zum Toggeln vorbereitet werden. Vervollständigt wird ein synchroner Zähler üblicherweise noch durch einen oder mehrere so genannte Freigabeeingänge, über die der Zähler gesteuert werden kann. Die Schaltung eines solchen synchronen Vorwärtzzählers zeigt Bild 9.30.

Für einen synchronen dualen Rückwärtzzähler lassen sich ebenfalls mit einer entsprechenden Zustandstabelle die Bedingungen, wann die einzelnen Flip-Flops umschalten dürfen, ableiten. Daraus ergibt sich dann eine Schaltung nach Bild 9.31. Im Gegensatz zur Schaltung des Vorwärtzzählers werden die Toggle-Eingänge der einzelnen Flip-Flops hier über die $\overline{Q}$ -Ausgänge freigegeben. Eine andere Möglichkeit für einen synchronen

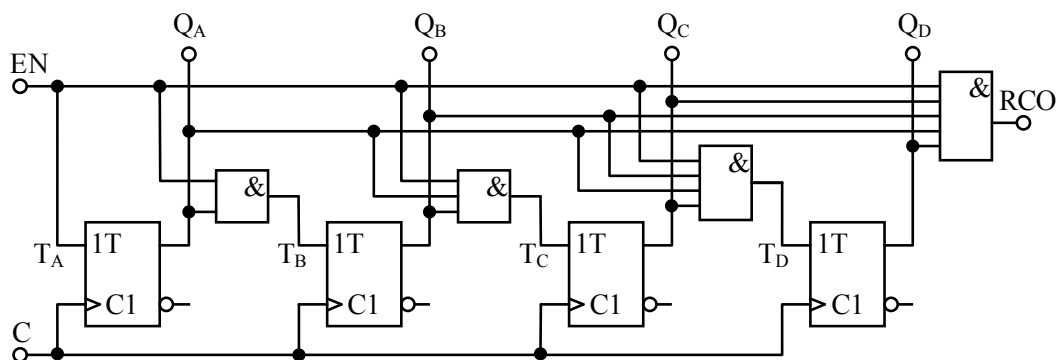


Bild 9.30: Synchroner Vorwärtszähler mit Freigabe-Eingang und Übertrag-Ausgang.

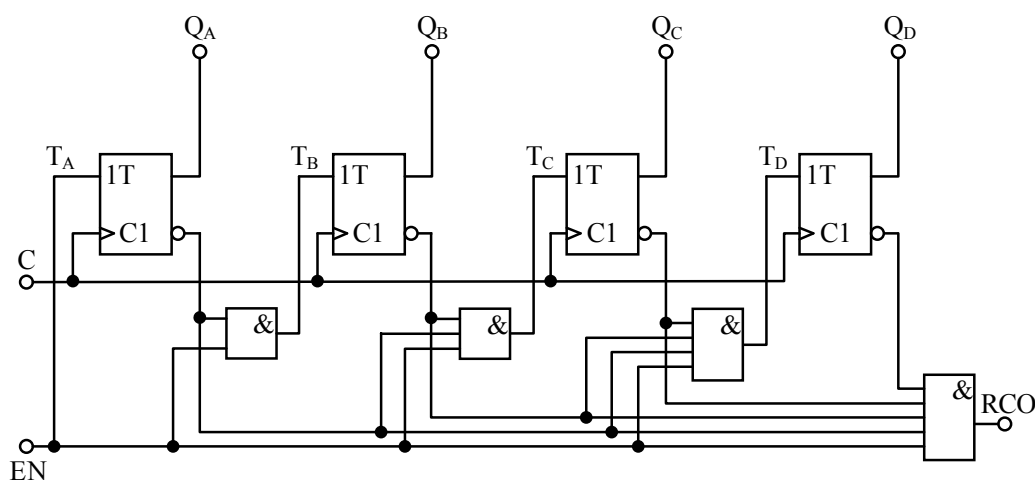


Bild 9.31: Synchroner Rückwärtszähler mit Freigabe-Eingang und Übertrag-Ausgang.

Rückwärtszähler ist, anstelle der Q -Ausgänge in Bild 9.30 die $\overline{Q}$ -Ausgänge als die Zähl-
ausgänge Q_A , Q_B , Q_C und Q_D zu verwenden. Eine Kombination eines Vorwärtszählers
nach Bild 9.30 und eines Rückwärtszählers nach Bild 9.31 ergibt unter Hinzufügen eines
zusätzlichen Steuereingangs einen kombinierten Vorwärts-Rückwärtszähler, wie er in
Bild 9.32 dargestellt ist. Eine logische 1 am EN-Eingang gibt die Funktion des Zählers
frei und eine 1 am $U/\overline{D}$ -Eingang gibt das Vorwärtszählen frei, während bei einer 0 an
diesen Eingang der Zähler rückwärts zählt. Anstelle eines OR-Gatters vor jedem T-
Eingang wurde im Schaltsymbol des T-Flip-Flops diese Verknüpfung in das Flip-Flop
integriert.

Synchroner BCD-Zähler

Bereits beim asynchronen BCD-Zähler haben wir festgestellt, dass ein Einfaches hin-
tereinander schalten von Flip-Flops ohne zusätzliche Logik nicht mehr möglich ist, da
für das Zählen einer Dekade andere Regeln als für rein binäres Zählen gelten. Deshalb
muss auch beim synchronen BCD-Zähler eine Überprüfung der Zustände am T-Eingang

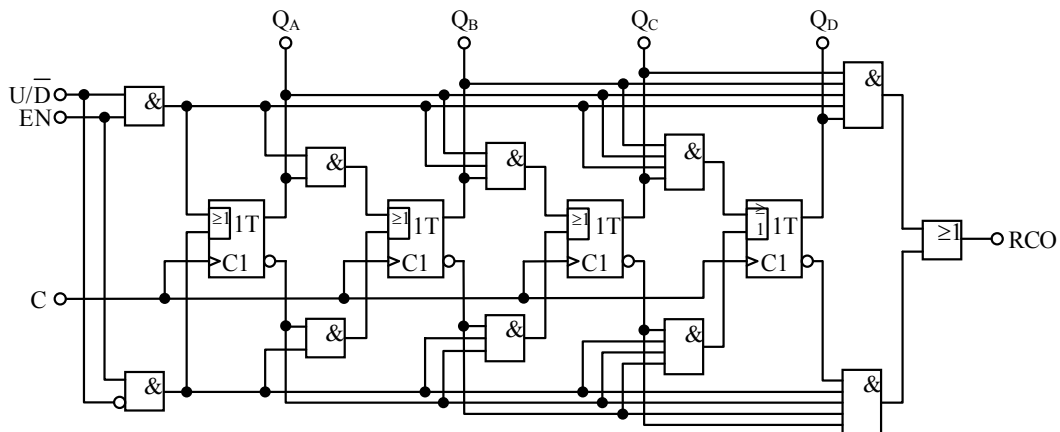


Bild 9.32: Synchroner Vorwärts-Rückwärtszähler.

stattfinden. Betrachten wir dazu noch einmal die Zustandstabelle 9.3:

Wir sehen, dass zunächst die logischen Pegel an den T-Eingängen der Flip-Flops identisch denen beim dualen Zähler sind. Entsprechend können auch die Verknüpfungen zu Vorbereitung der T-Eingänge der einzelnen Flip-Flops vorgenommen werden. Nach dem 8. Taktimpuls jedoch haben wir andere Bedingungen:

Takt (C)	Q_D	Q_C	Q_B	Q_A	Dezimalwert N	T_D	T_C	T_B	T_A
0,1	Q_{D-1}	Q_{C-1}	Q_{B-1}	Q_{A-1}	N_{-1}	x	x	x	x
⌋	0	0	0	0	0	0	0	0	1
⌋	0	0	0	1	1	0	0	1	1
⌋	0	0	1	0	2	0	0	0	1
⌋	0	0	1	1	3	0	1	1	1
⌋	0	1	0	0	4	0	0	0	1
⌋	0	1	0	1	5	0	0	1	1
⌋	0	1	1	0	6	0	0	0	1
⌋	0	1	1	1	7	1	1	1	1
⌋	1	0	0	0	8	0	0	0	1
⌋	1	0	0	1	9	1	0	0	1
⌋	0	0	0	0	10	0	0	0	1

Tabelle 9.3: Zustandstabelle eines BCD-Vorwärtszählers.

- beim zweiten Flip-Flop muss $T = 0$ bleiben, d.h. ein weiteres Umschalten des Ausgangs muss verhindert werden, aber

- beim vierten Flip-Flop muss $T = 1$ werden, damit sich nach dem nächsten Takt wieder der Ausgangszustand 0 0 0 0 einstellt.

Die erste Bedingung lässt sich wieder einfach lösen: wir nehmen wie beim asynchronen BCD-Zähler den Q-Ausgang des vierten Flip-Flops als zusätzliche Signal zur Steuerung des T-Eingangs des zweiten Flip-Flops. Damit ist gewährleistet, dass der T-Eingang in jedem Fall auf logisch 0 liegt. Jetzt muss nur noch dafür gesorgt werden, dass der T-Eingang des vierten Flip-Flops auf logisch 1 gelegt wird. Dazu können die Ausgänge Q_A und Q_D verwendet werden. Mit einem OR-Gatter werden die beiden Verknüpfungen für den T-Eingang des vierten Flip-Flops zusammengeführt. Bild 9.33 zeigt die Schaltung des beschriebenen synchronen BCD-Zählers.

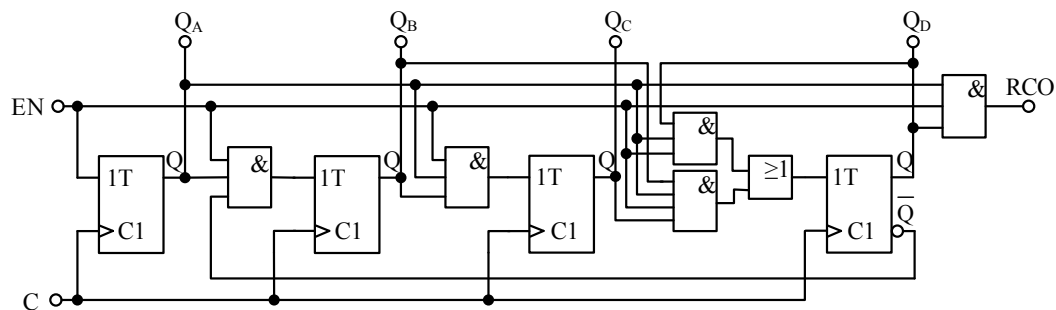


Bild 9.33: Schaltung eines synchronen BCD-Zählers für eine Dezimalstelle.

Synchrone Zähler mit mehreren Funktionen

Für eine Reihe von Anwendungen sind synchrone Zählerschaltungen notwendig, die vor Beginn eines Zählvorgangs synchron (mit dem Takt) auf einen Startwert gesetzt werden müssen und asynchron wieder auf 0 0 0 0 wieder zurückgesetzt werden sollen. Sie werden auch als programmierbare Zähler bezeichnet. Um einen solchen komplexen Zähler aufbauen zu können müssen erst einige grundlegende Überlegungen über den Aufbau der einzelnen Zählstufen (Flip-Flops) angestellt werden. Wie wir bei den bisherigen Zählern gesehen haben, werden fast ausschließlich JK- bzw. T-Flip-Flops eingesetzt. Zuerst sollen die notwendigen Funktionen in Worten beschrieben werden:

- Rücksetzen auf den Zustand 0 0 0 0, wenn $R = 0$ wird und unabhängig davon, welche logischen Pegel oder Signale an den anderen Eingängen anliegen
- Laden mit einem Startwert D wenn ein Ladesignal $L = 1$ anliegt und die ansteigende Flanke des Taktsignals eintrifft
- Zählen, wenn ein Eingang $T = 1$ gesetzt ist und das Flip-Flop nicht mit einem Startwert D geladen wird, also $L = 0$ ist.

a)

$\overline{R}$	C	L	T	D	J	K	Q
0	x	x	x	x	x	x	0
1	0	x	x	x	x	x	Q_{-1}
1	1	x	x	x	x	x	Q_{-1}
1	$\lceil$	1	x	0	0	1	0
1	$\lceil$	1	x	1	1	0	1
1	$\lceil$	0	0	x	0	0	Q_{-1}
1	$\lceil$	0	1	x	1	1	$\overline{Q_{-1}}$

b)

Bild 9.34: Wahrheitstabelle und Schaltsymbol eines Multifunktions-Flip-Flops.

Mit diesen Bedingungen lässt sich eine Wahrheitstabelle (Bild 9.34a) für die gewünschten Funktionen mit einem JK-Flip-Flop erstellen. Genau so einfach ist es auch, das zugehörige logische Symbol (Bild 9.34b) für ein solches Multifunktions-Flip-Flop zu erstellen. Deutlich aufwendiger ist es jedoch, die notwendige Logik zur Bereitstellung der logischen Pegel am J- und K-Eingang zu entwerfen.

Dies soll in Bild 9.35 versucht werden. Als Basis-Flip-Flop wird ein einflankengesteuertes JK-Flip-Flop mit einem Rücksetzeingang verwendet. Damit bleiben für die Bereitstellung der logischen Signale am J- und K-Eingang noch D, L und T. Aus der Wahrheitstabelle lässt sich ablesen: wenn $L = 1$ ist, ist $J = D$ und $K = D$. Dies kann durch AND-Gatter und einem Inverter erzeugt werden. Wenn $L = 0$ ist muss an J und K der logische Pegel des T-Eingangs anliegen. Dies bedeutet, dass aus beiden Bedingungen die Eingangssignale J und K gebildet werden müssen. Durch das Hinzufügen von OR-Gattern kann dies erreicht werden. Damit lässt sich eine mögliche schaltungs-technische Lösung nach Bild 9.35 zur Realisierung der Wahrheitstabelle nach Bild 9.34a aufbauen. Mit solchen multifunktionellen Flip-Flops lassen sich dann programmierbare Vorwärts-Rückwärts-Zähler nach Bild 9.34 aufbauen. Der Eingang EN steuert die Freigabe des Zählers ($1 = \text{Zählen}$, $0 = \text{Halt}$). Der Eingang $U/\overline{D}$ legt die Zählrichtung ($1 = \text{Up}$, $0 = \text{Down}$) fest. Diese Art von programmierbaren Zählern gibt es in einer Reihe von Ausführungen. Allen gemeinsam ist, dass sie sich sehr leicht zu Zählern höherer Bitbreite kaskadieren lassen, da aufgrund der Eigenschaften der synchronen Zähler immer nur die Gatterlaufzeit eines einzelnen Flip-Flops als Verzögerung zwischen dem Takteingang und allen Ausgängen des Zählers auftritt.

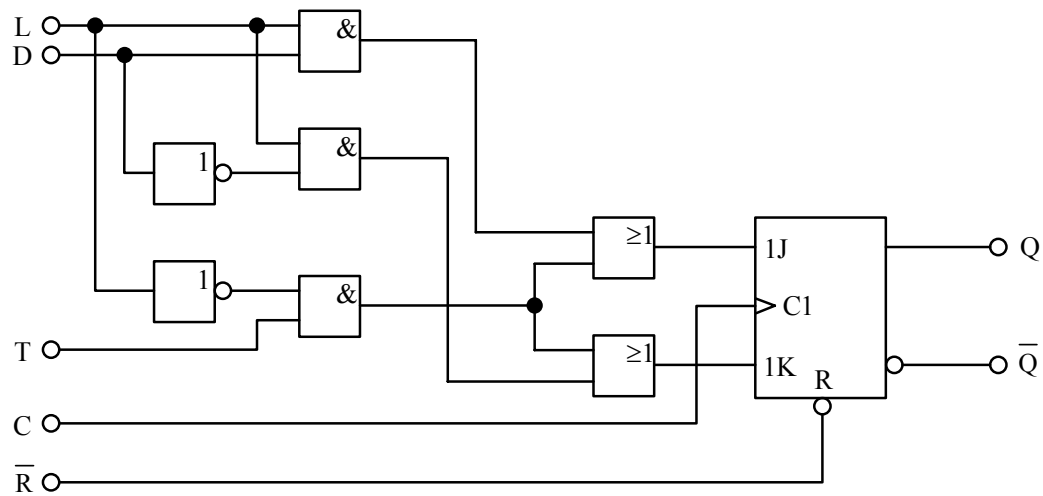


Bild 9.35: Möglicher Schaltungsaufbau des Multifunktions-Flip-Flops.

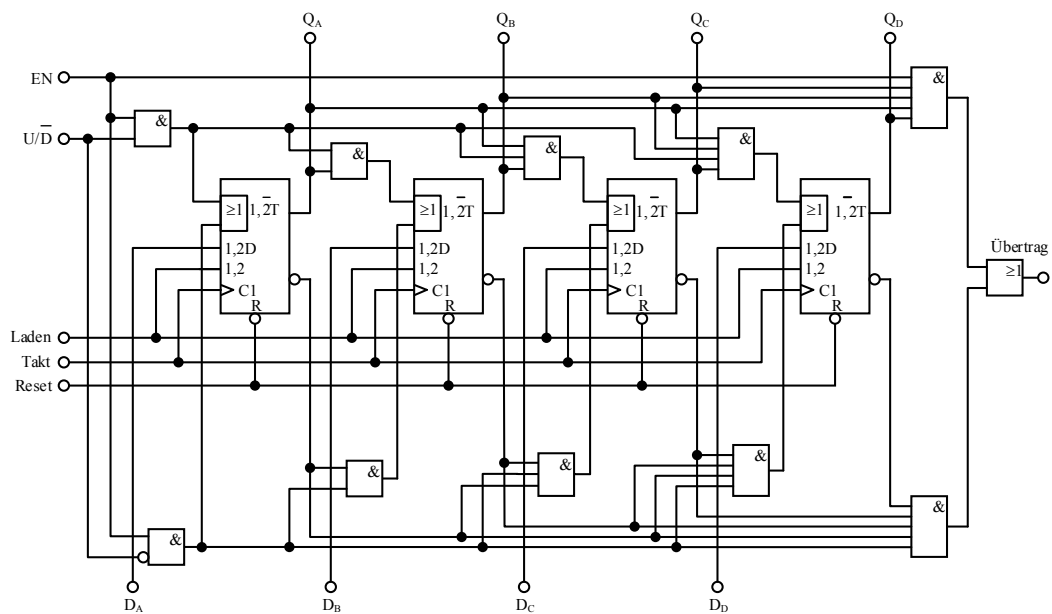


Bild 9.36: Schaltung eines programmierbaren synchronen dualen Vorwärts-Rückwärts-Zählers.

9.2.3 Zähler als Frequenzteiler

Frequenzteiler teilen die Frequenz eines rechteckförmigen Signals in einem bestimmten Verhältnis. Ein einzelnes Flip-Flop erzeugt eine Frequenzteilung im Verhältnis 2:1. Mit 2 Flip-Flops kann man ein Teilverhältnis von 4:1 erreichen. Man unterscheidet Frequenzteiler mit festem Teilverhältnis und Frequenzteiler, deren Teilverhältnis in einem gewissen Bereich einstellbar ist. Letztere werden auch programmierbare Frequenzteiler genannt.

9.2.3.1 Asynchrone Frequenzteiler mit festem Teilverhältnis

Jeder asynchrone Dualzähler eignet sich auch als Frequenzteiler mit einem festen Teilverhältnis. Betrachten wir als Beispiel die Schaltung und das Zeitablaufdiagramm eines 3-Bit Dual-Vorwärtszählers in Bild 9.37. Diese Schaltung basiert wiederum auf T-Flip-Flops, die aus JK-Flip-Flops aufgebaut wurde. Das erste Flip-Flop des Zählers halbiert die Frequenz des Eingangssignals E . Das zweite Flip-Flop halbiert die schon halbierte Frequenz ein weiteres Mal. Nochmals wird die Frequenz durch das dritte Flip-Flop halbiert. Ein 3-Bit Dual-Vorwärtszähler liefert also als Frequenzteiler an den Ausgängen Q_A , Q_B und Q_C die Teilverhältnisse 2:1, 4:1 und 8:1.

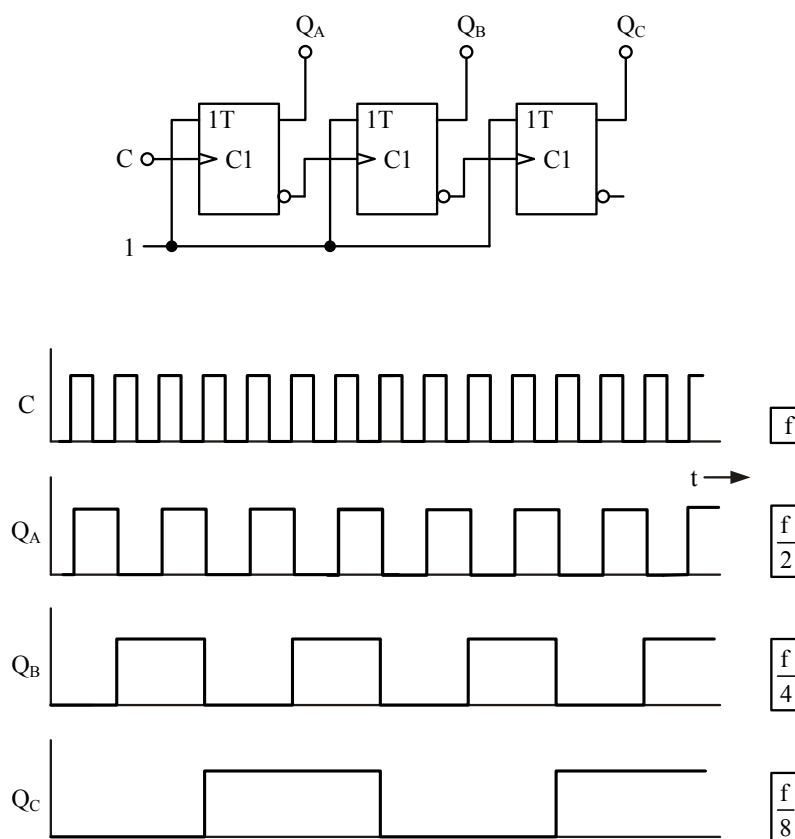


Bild 9.37: Asynchroner Vorwärtszähler als Frequenzteiler.

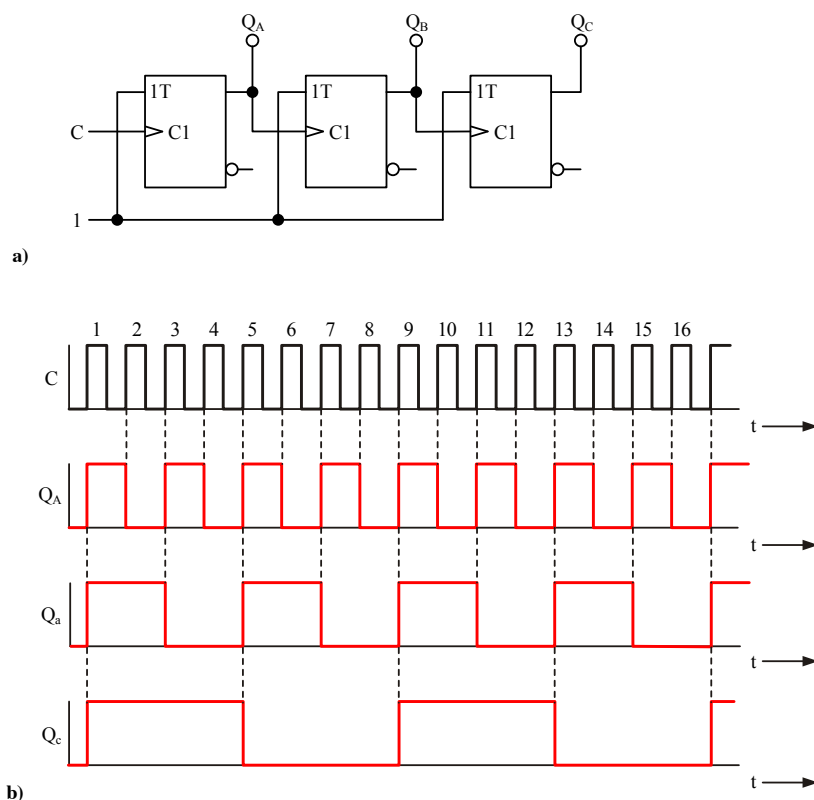


Bild 9.38: Asynchroner Rückwärtszähler als Frequenzteiler.

Dual-Rückwärtszähler sind ebenfalls als Frequenzteiler geeignet. Bild 9.38 zeigt die Schaltung und das Zeitablaufdiagramm eines solchen Frequenzteilers. Die geteilten Signale haben lediglich eine andere Phasenlage als bei dualen Vorwärtszählern. Geradzahlige Teilverhältnisse nach der 2-er Potenzreihe lassen sich also leicht erreichen. Jedes Flip-Flop teilt um den Faktor 2, d.h. es gilt die Gleichung $f_T = f_E/2^n$, wobei n die Anzahl der Flip-Flops, f_T die geteilte Frequenz und f_E die Eingangsfrequenz sind.

Um ungeradzahlige Teilverhältnisse zu erreichen, muss man Flip-Flops verwenden, die Rückstelleingänge haben. Ein Frequenzteiler mit dem Teilverhältnis 3:1 ist im Bild 9.39 dargestellt. Das Ausgangssignal Q_B hat ein anderes Impuls-Pausenverhältnis als das Eingangssignal E. Das ist für viele Anwendungsfälle ungünstig. Schaltet man aber ein weiteres Flip-Flop nach, ergibt sich wieder ein Impulspausenverhältnis von 1:1, wie es in Bild 9.40 gezeigt ist. Damit entsteht wiederum ein Teiler mit einem Teilverhältnis 6:1. In der Praxis werden häufig Frequenzteiler mit einem Teilverhältnis 10:1 benötigt. Ein solcher Teiler lässt sich einfach durch die Kombination eines 5:1 Teilers mit einem 2:1 Teiler realisieren. Bild 9.41 zeigt einen Frequenzteiler mit dem Verhältnis 10:1 und das zugehörige Zeitablaufdiagramm. Im linken Teil der Schaltung ist der 5:1 Teiler realisiert. Sein Ausgang ist mit einem weiteren 2:1 Frequenzteiler verbunden, der dann entsprechend ein weiteres Mal teilt und somit ein Gesamtteilerverhältnis 10:1 erzeugt. Das Teilverhältnis 5:1 wird dadurch erreicht, dass der Ausgang Q_A und Q_C über ein UND-

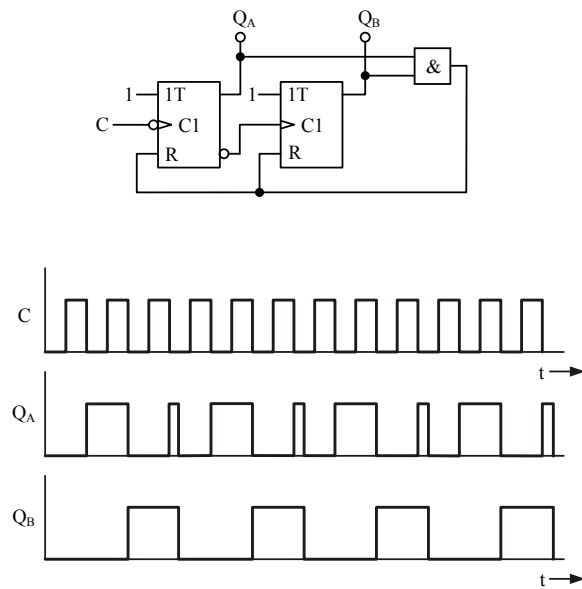


Bild 9.39: Asynchroner Frequenzteiler mit dem Teilerverhältnis 3:1.

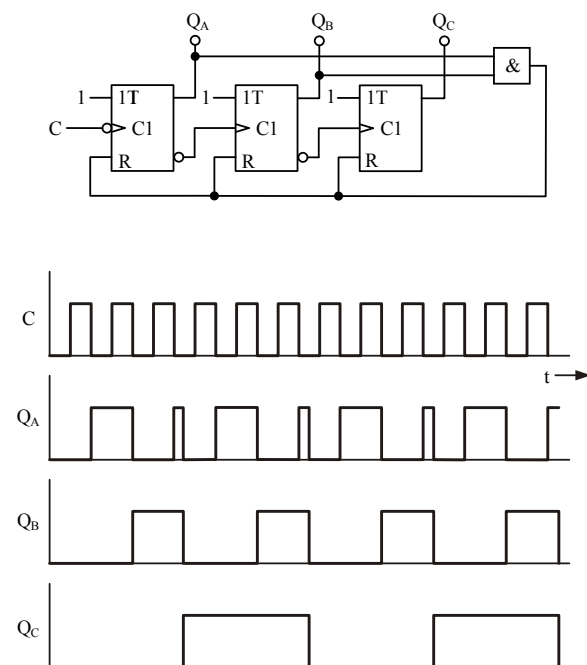


Bild 9.40: Frequenzteiler mit einem Teilerverhältnis 6:1.

Gatter mit den Rückstelleingängen der T-Flip-Flops verbunden ist und entsprechend nach dem 5. Zählimpuls zurückschaltet.

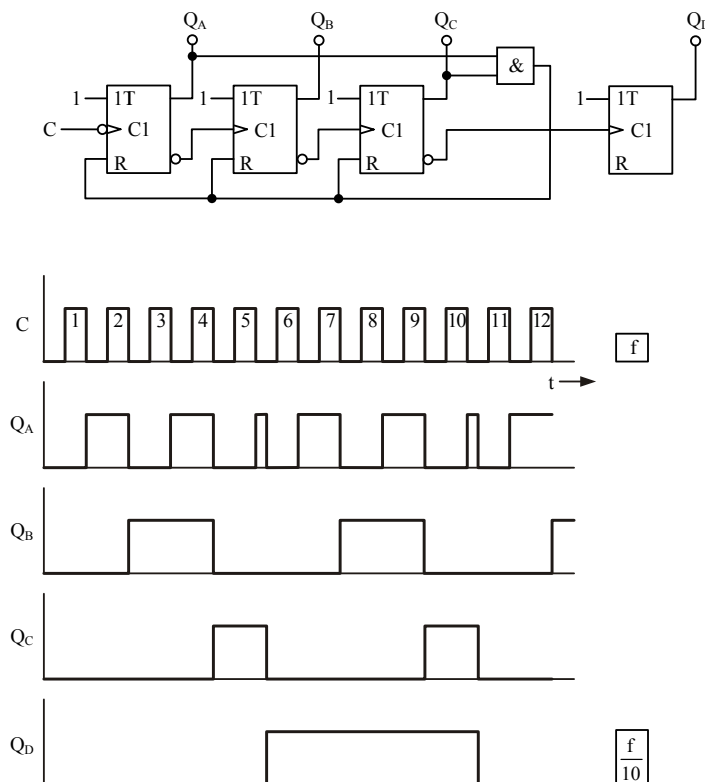


Bild 9.41: Frequenzteiler mit einem Teilverhältnis 10:1.

9.2.3.2 Synchrone Frequenzteiler mit festem Teilverhältnis

Jeder synchrone Dualzähler kann auch als Frequenzteiler mit festem Teilverhältnis arbeiten. Das gilt unter Einschränkungen nur für die Teilverhältnisse, die zur 2-er Potenzreihe gehören. Für andere Teilverhältnisse, insbesondere für ungerade, muss die Beschaltung der einzelnen Eingänge der JK-Flip-Flops geändert werden. Bild 9.42 zeigt die Schaltung und das Zeitablaufdiagramm eines synchron arbeitenden Frequenzteilers mit einem Teilverhältnis 3:1. Um eine eindeutige Funktion des Frequenzteilers zu erhalten, müssen die Gatterlaufzeiten zwischen Takteingang und Ausgang der Flip-Flops berücksichtigt werden. Ausgehend von einem Zustand $Q_A = 0$ und $Q_B = 0$ ergeben sich: $1J = 1$, $K = 1$ und $2J = 0$, $2K = 1$. Bei der ansteigenden Flanke des ersten Taktimpulses kann zwar das erste Flip-Flop umschalten, das zweite jedoch nicht, da $2J$ ja noch 0 ist. Bei der ansteigenden Flanke des zweiten Taktimpulses kann das erste Flip-Flop wieder toggeln und das zweite Flip-Flop wird gesetzt. Beim dritten Taktimpuls ist $1J = 1K = 0$, $2J = 0$ und $2K = 1$. Damit bleibt der Ausgang des ersten FF unverändert und das zweite FF wird zurückgesetzt. Ab dem vierten Taktimpuls wiederholt sich der gesamte beschriebene Vorgang von neuem.

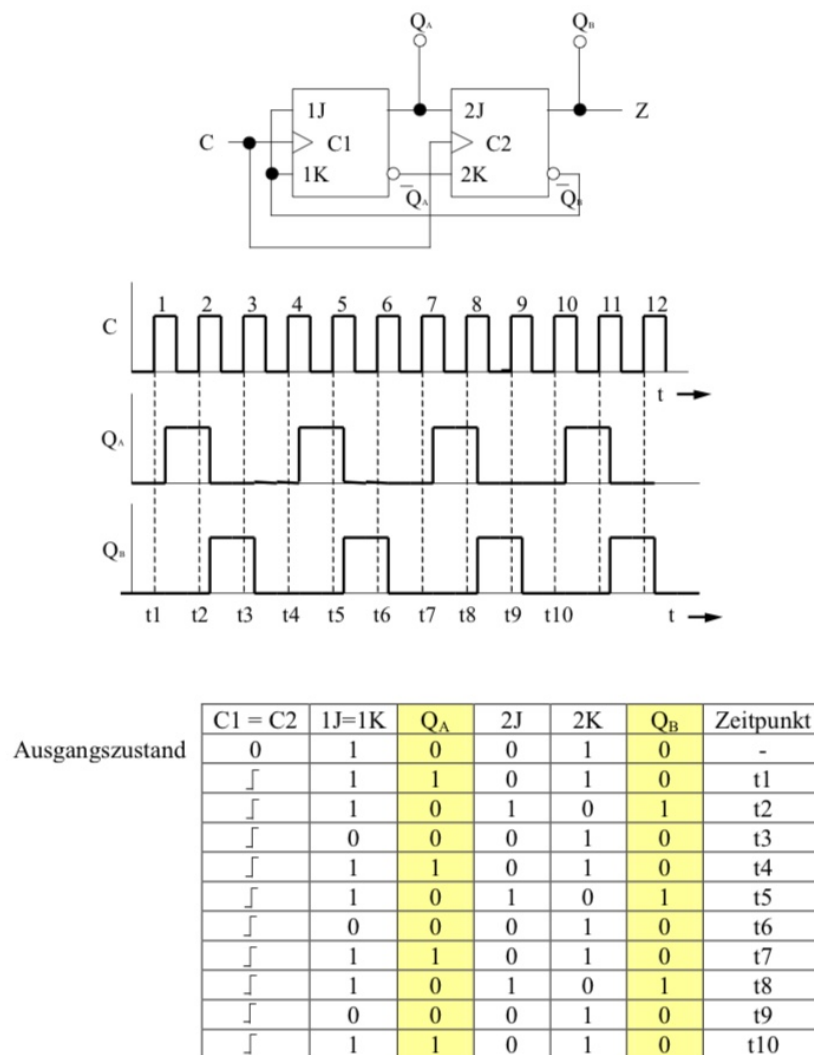
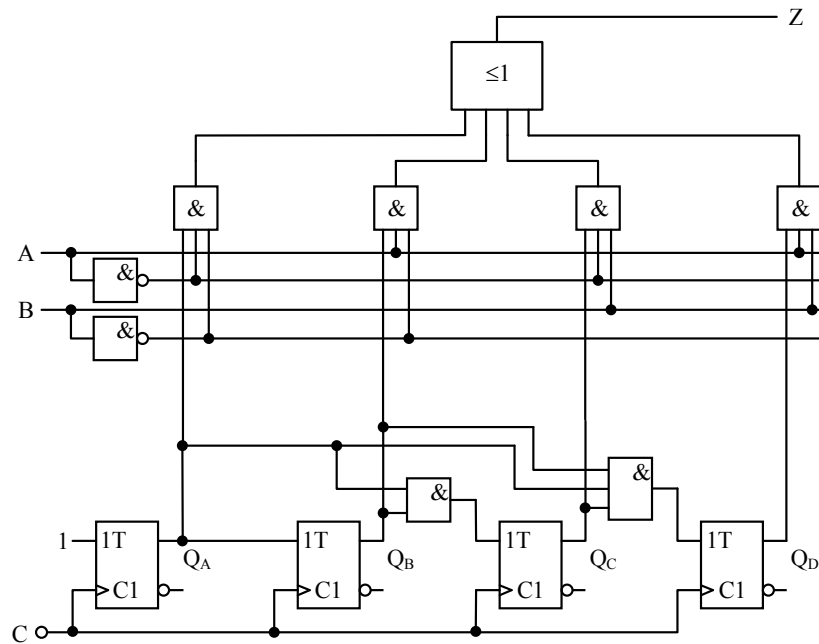


Bild 9.42: Synchroner Frequenzteiler mit einem Teilverhältnis 3:1.

9.2.3.3 Synchrone Frequenzteiler mit einstellbarem Teilverhältnis

Frequenzteiler mit einstellbarem Teilverhältnis sind im Prinzip synchrone Vorwärtszähler. Sie führen mehrere Frequenzteilungen durch. Das Signal mit der gewünschten Frequenzteilung wird über eine Auswahlschaltung auf den Ausgang gegeben. Ein Beispiel eines einstellbaren Frequenzteilers ist in Bild 9.43 gezeigt. Dieser Teiler teilt mit den Teilverhältnissen 2:1, 4:1, 8:1 und 16:1. Durch die Auswahlssignale A und B wird das gewünschte Signal auf den Ausgang Z geschaltet. Die Schaltung des Frequenzteilers kann darüber hinaus durch Umschaltungen verändert werden, so dass sich auch verschiedene ungerade Teilverhältnisse erreichen lassen.



B	A	Frequenz- teilverhältnis
0	0	2 : 1
0	1	4 : 1
1	0	8 : 1
1	1	16 : 1

Bild 9.43: Synchroner Frequenzteiler mit einstellbarem Teilverhältnis.

9.3 Schieberegister

Lernziele

- Kennenlernen und Verstehen des Aufbaus und der Wirkungsweise von Schieberegistern

Schieberegister sind Schaltungen, die eine Information taktgesteuert bitweise aufnehmen, sie eine gewisse Zeit speichern und dann wieder abgeben. Schieberegister werden aus Flip-Flops aufgebaut. Gut dazu eignen sich taktflankengesteuerte D-Flip-Flops, RS-Flip-Flops und JK-Flip-Flops. Hochwertige Schieberegister werden in der Regel mit JK-Master-Slave Flip-Flops aufgebaut. Diese Schaltungen stehen in der Regel als integrierte Schaltkreise zur Verfügung.

9.3.1 Schieberegister für serielle Ein- und Ausgabe

Ein einfaches Schieberegister mit 4 bit Speicherkapazität ist in Bild 9.44 dargestellt.

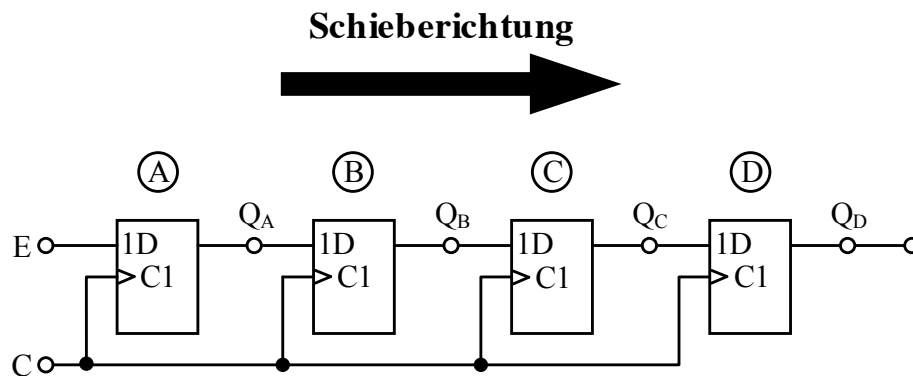


Bild 9.44: Schieberegister mit 4 bit für serielle Ein- und Ausgabe.

Es besteht aus vier D-Flip-Flops, die mit der positiven Taktflanke schalten. Liegt ein 1-Pegel am Eingang E und ändert sich das Taktsignal von 0 auf 1, so wird das Flip-Flop A gesetzt. An seinem Ausgang Q_A erscheint eine 1. Wird dann an den Eingang 0-Pegel belegt, so wird mit der zweiten ansteigenden Taktflanke das Flip-Flop A zurückgesetzt und das Flip-Flop B gesetzt. Die 1 erscheint jetzt am Ausgang Q_B . Mit der dritten ansteigenden Taktflanke wird das Flip-Flop B zurückgesetzt und das Flip-Flop C gesetzt. Der Ausgang Q_C wird jetzt 1. Mit der vierten ansteigenden Taktflanke wird das Flip-Flop C zurückgesetzt und das Flip-Flop D gesetzt ($Q_D = 1$). Der 1-Pegel, der zu Beginn am Eingang E anlag, wurde auf diese Weise Takt für Takt von Flip-Flop zu Flip-Flop weitergeschoben. Er liegt jetzt am Ausgang des Schieberegisters Q an. Mit der ansteigenden Flanke des fünften Taktes wird das Flip-Flop D ebenfalls zurückgesetzt. Man sagt, das Schieberegister sei jetzt "leer". Es enthält keine Informationen mehr.

Takt Nr. n	E	Zustände nach Takt Nr. n			
		Q_A	Q_B	Q_C	$Q_D = Q$
	1	0	0	0	0
1	0	1	0	0	0
2	0	0	1	0	0
3	0	0	0	1	0
4	0	0	0	0	1
5	0	0	0	0	0

Bild 9.45: Funktionstabelle des 4 bit Schieberegisters.

In Bild 9.45 sind die einzelnen Schiebeschritte in einer Funktionstabelle zusammengefasst. Das dazugehörige Zeitdiagramm ist in Bild 9.46 gezeigt. Man spricht von einer seriellen Dateneingabe, weil die Signale am Eingang E in das Schieberegister zeitlich nacheinander aufgenommen werden. Nach der Aufnahme der Information können die Taktsignale gesperrt werden. Die Information wird dann gespeichert und zwar solange,

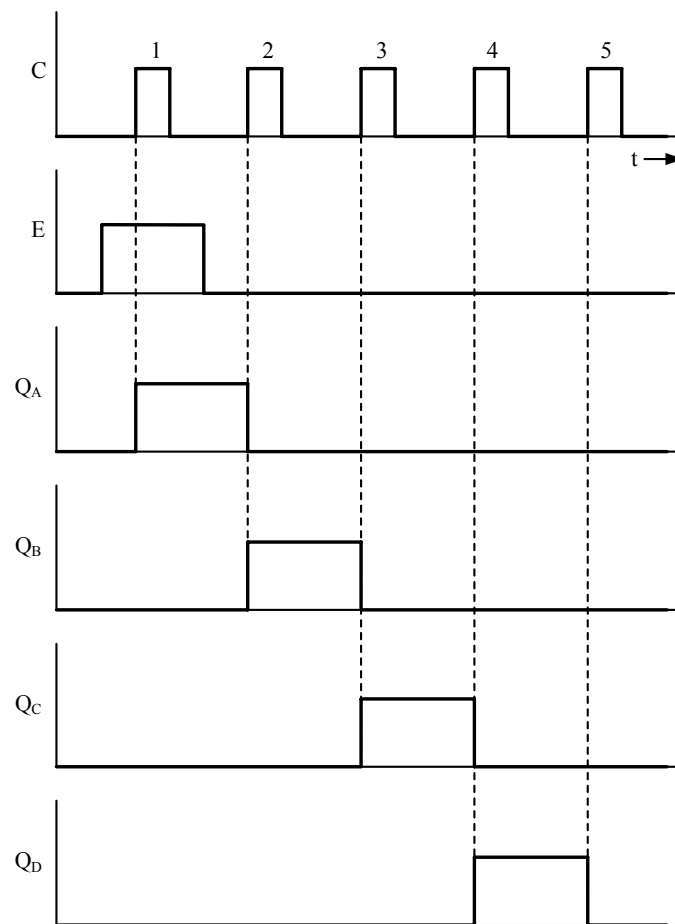


Bild 9.46: Zeitdiagramm für das Schieberegister nach Bild 9.44.

wie die Taktsignale gesperrt sind. Erfolgt eine Freigabe des Taktes, werden die Informationen bit nach bit an den Ausgang Q weitergegeben. Deshalb spricht man von einer seriellen Datenausgabe. Als Beispiel soll nun noch das Bitmuster 0101 in das Schieberegister nach Bild 9.44 eingegeben werden. Hierzu werden 4 Takte benötigt. Vor dem 1. Takt muss der Inhalt des 1. bit am Eingang E anliegen (1-Pegel). Vor dem 2. Takt muss der Inhalt des 2. bit am Eingang E anliegen. Das ist in diesem Fall ein 0-Pegel. Vor dem 3. Takt muss der Inhalt des 3. bit am Eingang liegen (1-Pegel). Vor dem 4. Takt muss der Inhalt des 4. bit am Eingang E anliegen (0-Pegel). Nach dem 4. Takt ist das Bitmuster 0101 in das Register eingeschrieben. Das Zeitablaufdiagramm in Bild 9.47 gibt diesen Vorgang wider. Wenn jetzt die Taktsignale gesperrt werden, kann die Information beliebig lange im Schieberegister gespeichert werden. Um die Information aus dem Schieberegister Ausgang Q auszulesen, werden weitere 4 Takte benötigt. Das 1. bit liegt ja bereits am Ausgang Q. Das 2. bit nach ist dem 5. Takt verfügbar. Das 3. bit liegt nach dem 6. Takt und das 4. bit nach dem 7. Takt am Ausgang Q. Nach dem 8. Takt ist das Schieberegister leer.

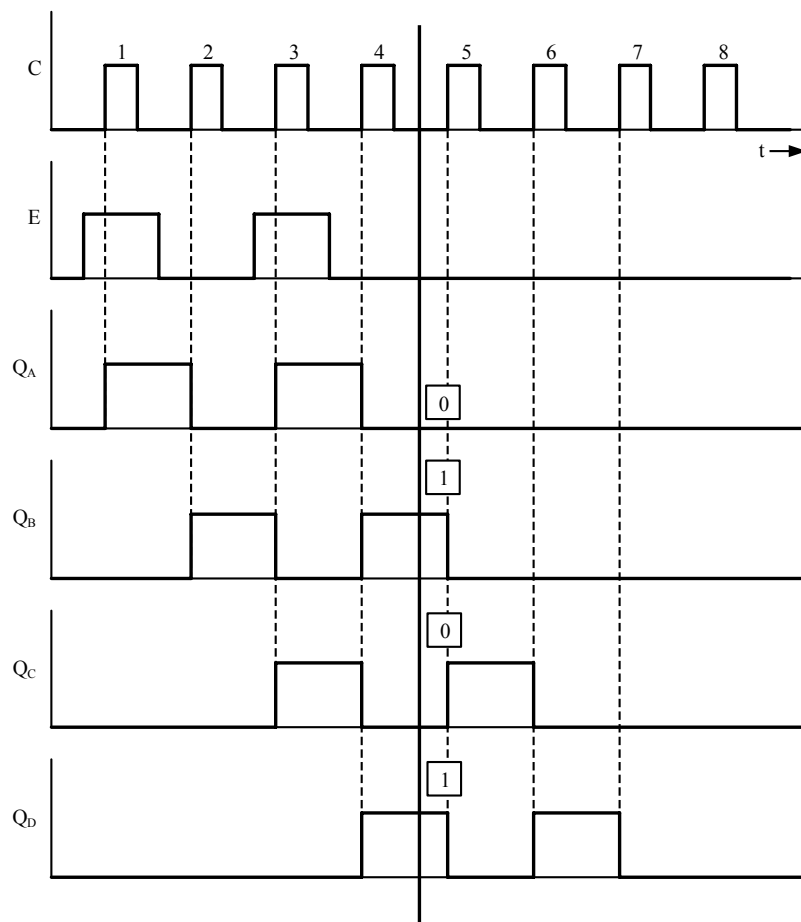


Bild 9.47: Zeitdiagramm für das Schieberegister nach Bild 9.44 mit der Bitfolge 0101.

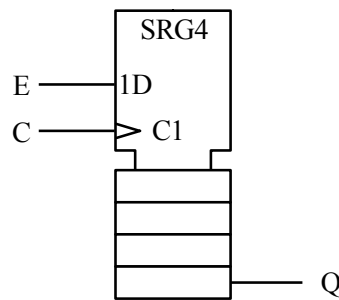


Bild 9.48: Schaltsymbol für ein Schieberegister mit 4 bit für serielle Ein- und Ausgabe.

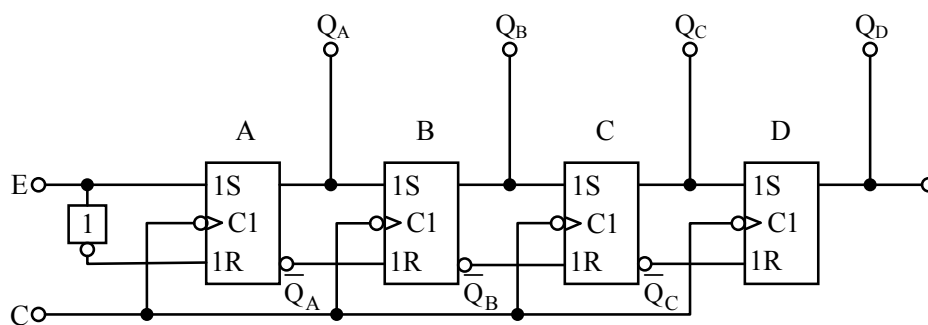


Bild 9.49: Schieberegister mit seriell und parallelem Ausgang.

Bild 9.48 zeigt das Schaltzeichen des beschriebenen 4 bit Schieberegisters, das mit D-Flip-Flop aufgebaut ist und mit serieller Ein- und Ausgabe arbeitet.

9.3.2 Schieberegister mit Parallelausgabe

Schieberegister haben immer die Möglichkeit der seriellen Dateneingabe und der seriellen Datenausgabe. Ohne diese Möglichkeit könnte eine Schaltung nicht als Schieberegister bezeichnet werden. Innerhalb des Registers stehen die Daten an den Ausgängen der einzelnen Flip-Flops parallel zur Verfügung. Bei einem Schieberegister mit Parallelausgabe werden zusätzlich diese innerhalb des Registers vorliegenden gespeicherten Daten auch parallel ausgegeben. Das Schieberegister nach Bild 9.49 hat die Möglichkeit der Parallelausgabe. Die Q-Ausgänge der Flip-Flops werden zusätzlich an Anschlusspunkte geführt. Dort sind dann alle 4 bit parallel verfügbar. Das Schieberegister nach Bild 9.49 wurde aus taktflankengesteuerten RS-Flip-Flops aufgebaut. Ein Zurücksetzen kann bei diesen Flip-Flops bekanntlich nur erfolgen, wenn am R-Eingang eine logische 1 und am S-Eingang eine logische 0 anliegt und die schaltende Taktflanke ankommt. Eine logische 0 am R-Eingang löst kein Kippen aus. Damit kein unerlaubter Eingangszustand $R = S = 1$ auftreten kann, wird das Eingangssignal direkt an den S-Eingang und über einen Inverter an den R-Eingang des 1. Flip-Flops geführt. Die nachfolgenden Flip-Flops werden dann von den $Q(=S)$ - und $Q(=R)$ -Ausgängen des jeweils vorausgehenden Flip-Flops angesteuert. Liegt am Eingang E ein 0-Pegel an, so liegt am Eingang

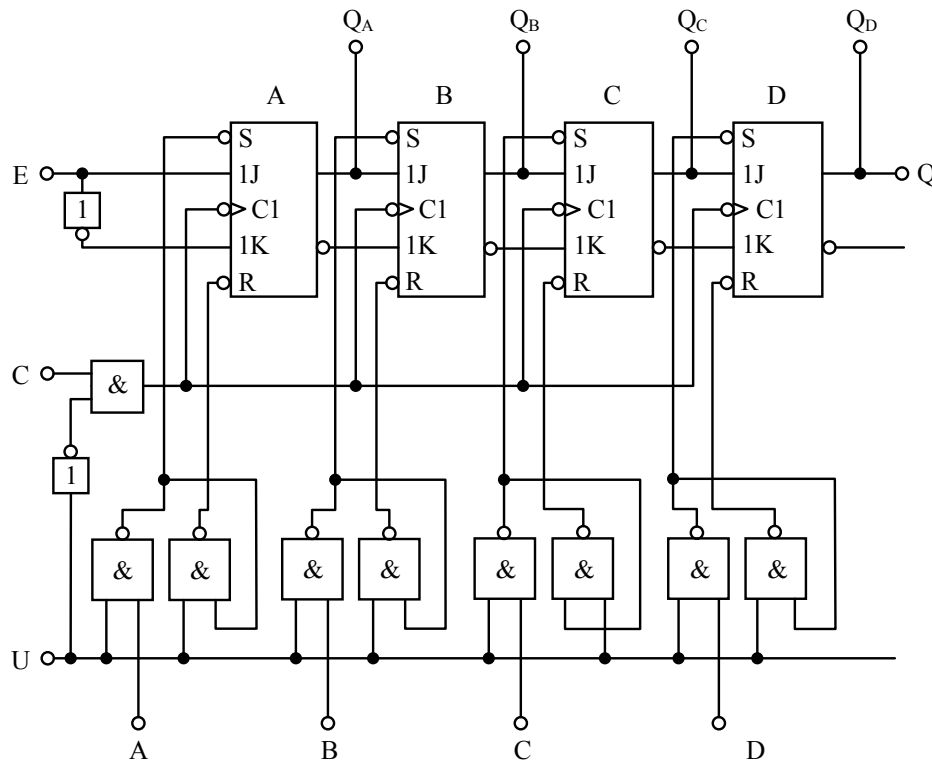


Bild 9.50: Schieberegister mit Parallelausgabe und Paralleleingabe.

R des 1. Flip-Flop ein 1-Pegel an, und das Flip-Flop wird mit der nächsten fallenden Flanke des Taktsignals zurückgesetzt. Die im Schieberegister gespeicherte Information kann taktunabhängig an den Ausgängen Q_A , Q_B , Q_C und Q_D abgenommen werden. Während dieser Parallelausgabe darf das Schieberegister keine weiteren Schiebetakte erhalten. Sonst würde die parallel ausgegebene Information verfälscht. Es darf also nicht gleichzeitig eine serielle und eine parallele Datenausgabe erfolgen. Ebenfalls darf nicht gleichzeitig eine serielle Dateneingabe und eine parallele Datenausgabe stattfinden.

9.3.3 Schieberegister mit Parallelausgabe und Paralleleingabe

Für viele Anwendungsfälle ist es günstig, neben der seriellen Dateneingabe die Möglichkeit zu haben, dem Schieberegister Daten parallel einzugeben. Diese Paralleleingabe kann taktabhängig oder taktunabhängig erfolgen. Das Schieberegister nach Bild 9.50 bietet neben der Möglichkeit der Parallelausgabe auch die Möglichkeit der Paralleleingabe. Parallelausgabe und Paralleleingabe sind taktunabhängig. Das Schieberegister ist mit JK-Flip-Flops aufgebaut, die taktunabhängige Stell- und Rückstelleingänge haben. Die Dateneingänge für die Paralleleingabe sind A, B, C und D. Die Paralleleingabe und serielle Ein- und Ausgabe sind gegeneinander verriegelt. Liegt am Umschalteingang ein 0-Pegel, so ist der Takt freigegeben. Das Schieberegister kann seriell arbeiten. Bei $U = 1$ ist eine Paralleleingabe möglich. Das Taktsignal ist gesperrt. Ein Schieberegister mit taktabhängiger Paralleleingabe wird in Bild 9.51 gezeigt. Die Eingänge J und K ei-

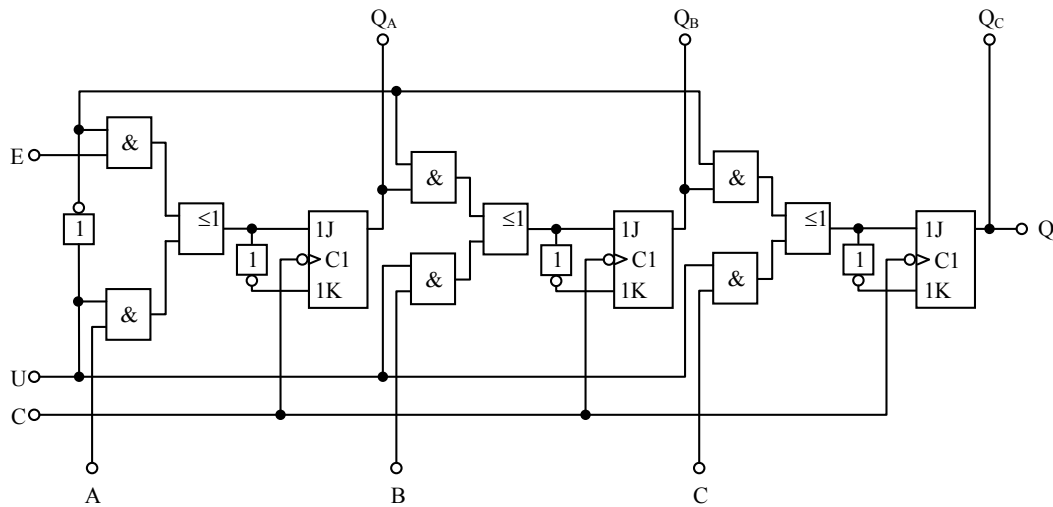


Bild 9.51: Schieberegister mit taktabhängiger Paralleleingabe.

nes jeden Flip-Flops sind umschaltbar. Bei $U = 0$ werden die Flip-Flop-Eingänge seriell mit Signalen versorgt. Bei $U = 1$ erhalten die Flip-Flop-Eingänge ihre Signale von den Dateneingängen für die Paralleleingabe (A, B, C, D). Das Setzen oder Rücksetzen der Flip-Flops erfolgt mit dem Takt. Im vorliegenden Fall erfolgt das Rücksetzen also mit der abfallenden Taktflanke. Typische Anwendungen solcher Schieberegister sind serielle Schnittstellen in PCs oder Ausgangsregister für synchrone DRAM-Speicher.

9.3.4 Ringregister

Ein Ringregister ist ein Schieberegister, dessen Ausgang mit dem Eingang verbunden ist. Bei einem Ringregister können die Informationen im Ring geschoben werden. Sie laufen also im Kreis herum. Deshalb wird auch ein solches Register häufig als Umlaufregister bezeichnet. Der prinzipielle Aufbau eines Ringregisters ist in Bild 9.52 gezeigt. Die Informationen können dabei seriell oder parallel in ein Ringregister eingegeben werden. Sie können ebenfalls seriell oder parallel ausgegeben werden. Bild 9.53 zeigt ein Ringregister mit serieller Dateneingabe und wahlweise serieller oder paralleler Datenausgabe. Bei $U = 1$ ist das Register als Ringregister geschaltet. Die Ausgangssignale werden vom Eingang aufgenommen. Bei $U = 0$ ist eine Dateneingabe über den seriellen Eingang E möglich. Eine serielle Datenausgabe ist über den Ausgang Q möglich, wenn $K = 1$ ist. Über R kann das Register taktunabhängig mit einem 0-Pegel zurückgesetzt werden, wobei die darin enthaltene Information gelöscht wird.

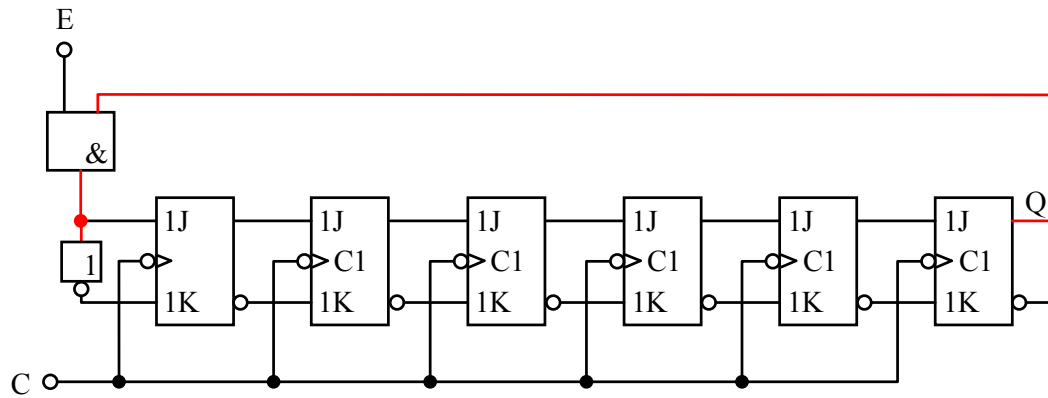


Bild 9.52: Prinzipieller Aufbau eines Ringregisters.

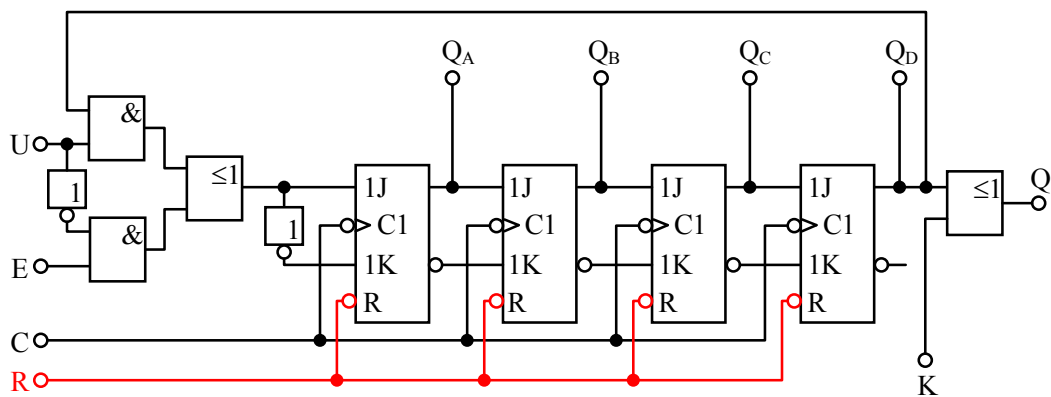


Bild 9.53: Aufbau eines Ringregisters mit wählbarer Funktion.

10 Decoder (Codewandler, Dekodierer), Multiplexer

Lernziele

- Kennenlernen und Verstehen des Aufbaus und der Wirkungsweise von Multiplexer- und Codewandlerschaltungen

Decoder sind integrierte Bausteine, die üblicherweise binäre Codes in eine andere Form umwandeln. Diese Form kann ein weiterer binärer Code oder aber ein ganz spezieller, auf eine bestimmte Anwendung zugeschnittener Code sein.

10.1 n bit binär zu 1-aus-n Decoder

1-aus-n Decoder sind Bausteine, die aus einer binären Eingangsinformation mit n Stellen eine einzelne Ausgangsinformation aus 2^n möglichen Werten aktivieren. Zur Verdeutlichung der Aufgabe ist es günstig, zuerst eine Wahrheitstabelle zu erstellen, wie sie in Bild 10.1a für einen 1-aus-4 Decoder gezeigt wird. Abhängig von der binären Information an den Eingängen $A_0 (= 2^0)$ und $A_1 (= 2^1)$ wird immer nur ein einzelner Ausgang aktiv (1). Alle anderen Ausgänge bleiben inaktiv (0). Eine einfache Schaltung zur Realisierung des Decoders ist in Bild 10.1b dargestellt. Wird anstelle des logischen 1-Pegels ein logischer 0-Pegel als aktives Signal benötigt, müssen die AND-Gatter einfach durch NAND-Gatter ersetzt werden.

10.2 BCD zu 7-Segment Decoder

Einen BCD zu 7-Segment Decoder haben wir bereits bei den Frequenzteilern als Funktionsblock eingesetzt. Damit soll die binäre Ausgangsinformation eines BCD-Zählers (eine Dezimalstelle) in eine Form zur Darstellung der Ziffer mit einer 7-Segment-Anzeige umgeformt werden. In Bild 10.2a ist ein solches Anzeigeelement mit den 7-Segmenten und deren Bezeichnung und daneben als Beispiel die Darstellung der Ziffer 5 gezeigt. Die Wahrheitstabelle (Bild 10.2b) zeigt, welche Segmente zur Darstellung der einzelnen Ziffern eingeschaltet werden müssen.

10.3 Prioritätsdecoder

Ein Prioritätsdecoder dient dazu, eine 1-aus-n Information in eine binäre Ziffer umzuwandeln, d.h. er führt die umgekehrte Aufgabe wie nach Bild 10.1 durch. An seinen Ausgängen tritt eine Dualzahl auf, die der höchstwertigen Stelle der Eingangsziffer entspricht.

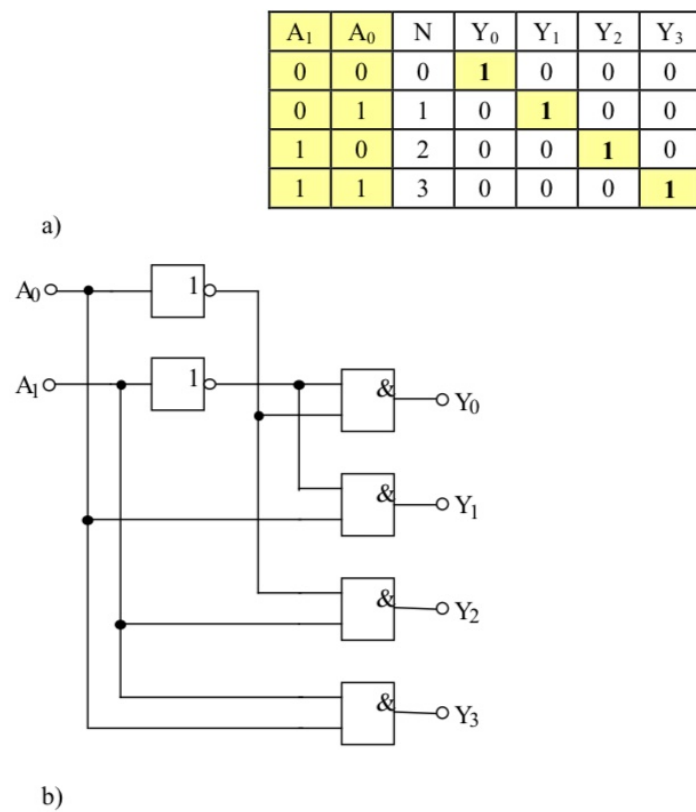
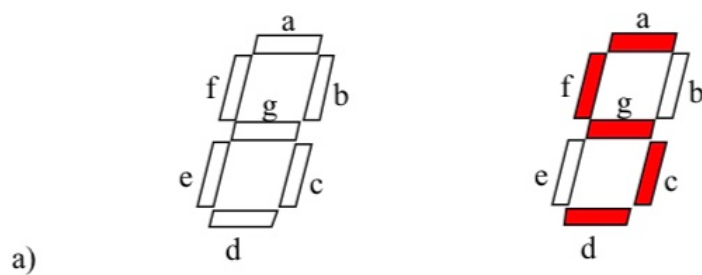


Bild 10.1: Wahrheitstabelle und Schaltungsaufbau eines 1-aus-4 Decoders.



Q_D	Q_C	Q_B	Q_A	Ziffer	a	b	c	d	e	f	g
0	0	0	0	0	1	1	1	1	1	1	0
0	0	0	1	1	0	1	1	0	0	0	0
0	0	1	0	2	1	1	0	1	1	0	1
0	0	1	1	3	1	1	1	1	0	0	1
0	1	0	0	4	0	1	1	0	0	1	1
0	1	0	1	5	1	0	1	1	0	1	1
0	1	1	0	6	1	0	1	1	1	1	1
0	1	1	1	7	1	1	1	0	0	0	0
1	0	0	0	8	1	1	1	1	1	1	1
1	0	0	1	9	1	1	1	1	0	1	1

b)

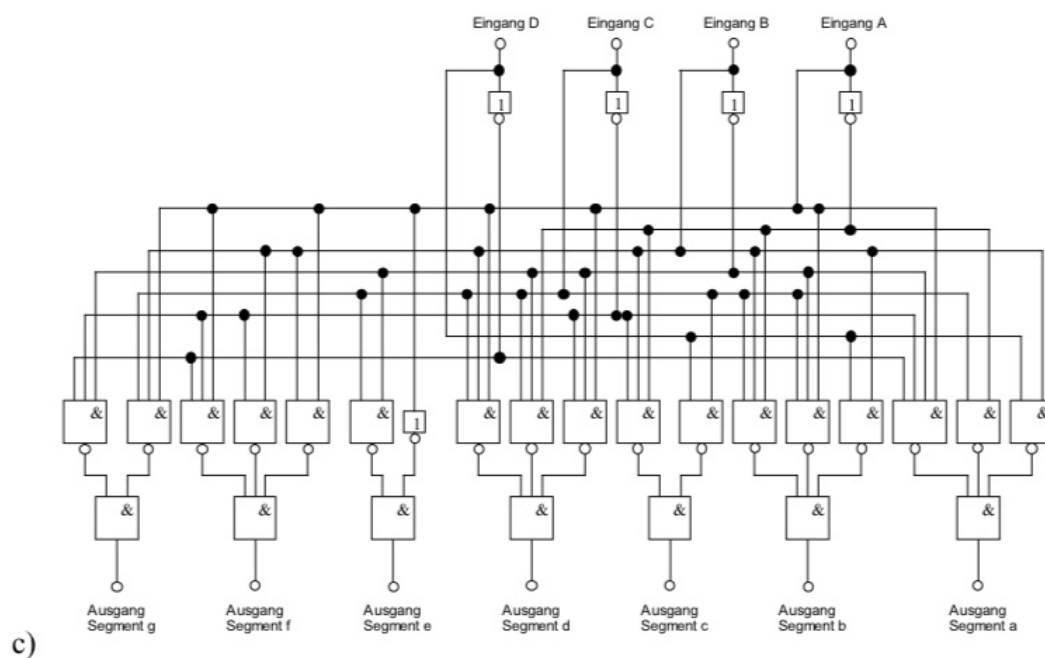


Bild 10.2: BCD zu 7-Segment Decoder.

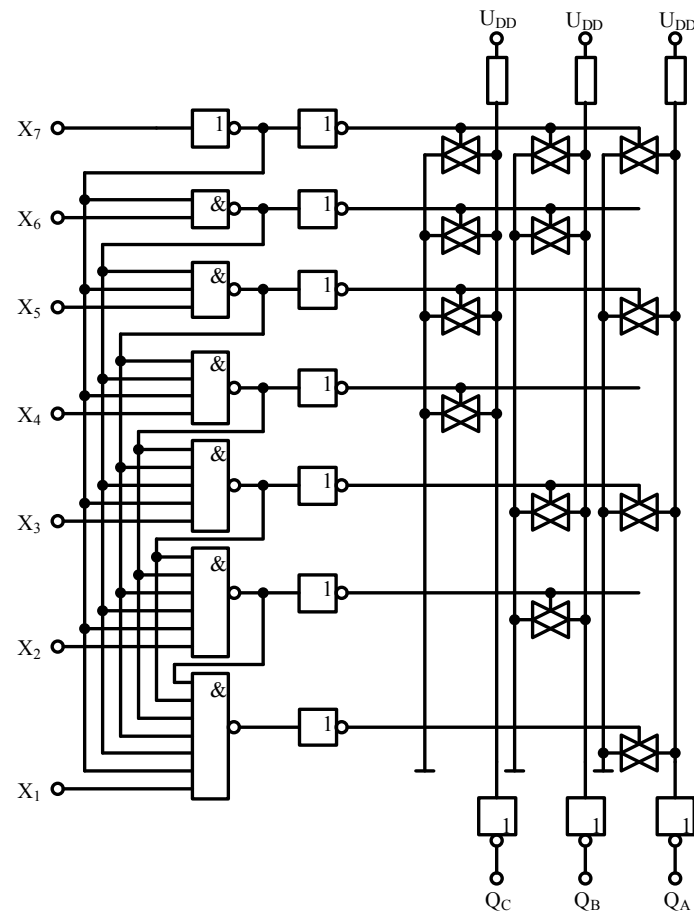


Bild 10.3: Prioritätsdecoder in CMOS-Technik.

X ₇	X ₆	X ₅	X ₄	X ₃	X ₂	X ₁	N	Q _C	Q _B	Q _A
0	0	0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	1	1	0	0	1
0	0	0	0	0	1	x	2	0	1	0
0	0	0	0	1	x	x	3	0	1	1
0	0	0	1	x	x	x	4	1	0	0
0	0	1	x	x	x	x	5	1	0	1
0	1	x	x	x	x	x	6	1	1	0
1	x	x	x	x	x	x	7	1	1	1

Bild 10.4: Wahrheitstabelle für einen Prioritätsdecoder.

Liegen gleichzeitig noch logische H-Pegel an niederwertigeren Stellen an, so werden diese ignoriert. Ein möglicher schaltungstechnischer Aufbau in CMOS-Technik ist in Bild 10.3 angegeben. Auch hier ist es von Vorteil, die Wahrheitstabelle der Schaltung (Bild 10.4) zu erstellen. Man erkennt, dass für den Ausgangszustand 0 0 0 an keinem

A ₁	A ₀	D ₀	D ₁	D ₂	D ₃	Y
0	0	0,1	X	X	X	D ₀
0	1	X	0,1	X	X	D ₁
1	0	X	X	0,1	X	D ₂
1	1	X	X	X	0,1	D ₃

Bild 10.5: Wahrheitstabelle eines Multiplexers mit vier Eingängen.

der Eingänge X_1 bis X_7 eine logische 1 anliegen darf. Erkennbar ist auch, dass immer die höchstwertige Ziffer für den binären Ausgangscode sorgt, egal welche logische Information an den niederwertigeren Eingängen anliegt. Daher auch die Bezeichnung "Prioritätsdecoder".

10.4 Multiplexer, Demultiplexer

Multiplexer und Demultiplexer sind integrierte Schaltungen, deren Funktion der eines mechanischen Drehschalters entspricht. Abhängig von der Stellung des Schalters wird einer von mehreren Eingängen auf einen Ausgang geschaltet. Ein Demultiplexer arbeitet gerade in umgekehrter Richtung. Man unterscheidet zwei Arten von Multiplexern:

- rein digitale Multiplexer (Bild 10.6) und
- sogenannte Analog-Multiplexer/Demultiplexer (Bild 10.7)

Zunächst soll aber auch hier eine Wahrheitstabelle (Bild 10.5) erstellt werden. Über die Adresseingänge A_0 und A_1 soll geschaltet werden, dass der logische Pegel (0 oder 1), der an einem von vier Dateneingängen D_0 bis D_3 anliegt, am Ausgang Y bereitstehen soll. Ein rein digitaler 4-zu-1 Multiplexer kann aus einem 1-aus-4 Decoder nach Bild 10.1 und Hinzufügen weiterer Eingänge für die Signale D_0 bis D_3 aufgebaut werden. Eine Verknüpfung der vier Ausgänge des Decoders über ein OR-Gatter zu einem einzelnen Ausgang erzeugt dann die gewünschte Funktion. Die Eingänge A_0 und A_1 stellen eine Adresse dar, durch die einer der vier Eingänge D_0 bis D_3 ausgewählt wird. Am Ausgang Y ist dann die Information des ausgewählten Eingangs nach folgender Boole'scher Gleichung abzulesen:

$$Y = \overline{A_0} \cdot \overline{A_1} \cdot D_0 + A_0 \cdot \overline{A_1} \cdot D_1 + \overline{A_0} \cdot A_1 \cdot D_2 + A_0 \cdot A_1 \cdot D_3$$

Mit einem Demultiplexer wird eine Eingangsinformation D entsprechend der binären Adresse an den Adresseingängen A_0 und A_1 an einen der Ausgänge Y_1 bis Y_3 geleitet. Die entsprechende Schaltung für einen rein digitalen Demultiplexer ist in Bild 10.7 dargestellt. Die Basis ist auch hier wieder ein 1-aus-4-Dekoder.

Das Prinzip einer Analog-Multiplexer- / Demultiplexer-Schaltung ist in Bild 10.8a gezeigt. Die Ausgänge eines 1-aus-4-Dekoders nach Bild 10.1 steuern vier Schalter an,

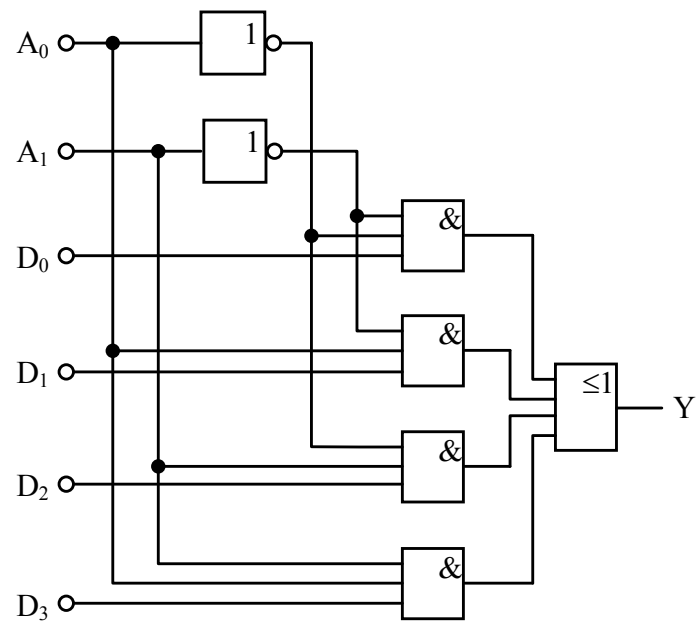


Bild 10.6: Schaltung eines digitalen Multiplexers.

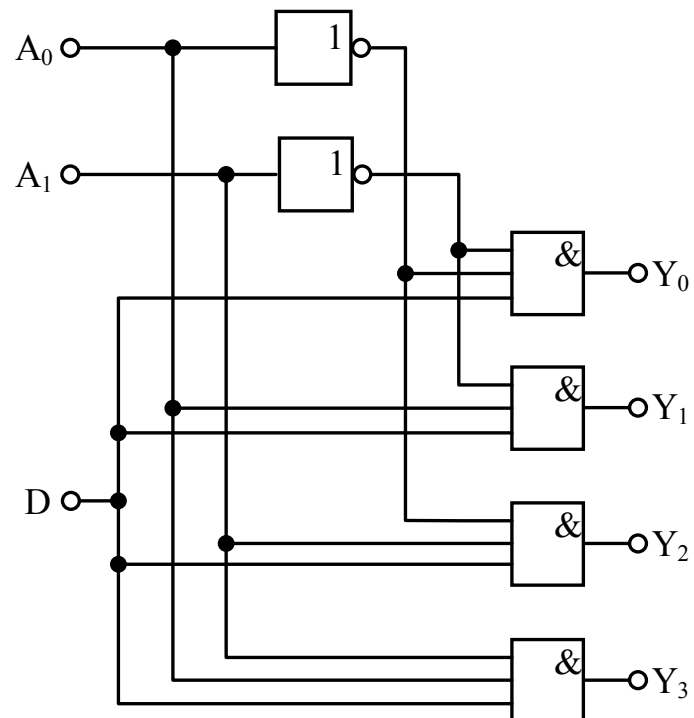


Bild 10.7: Schaltung eines digitalen Demultiplexers.

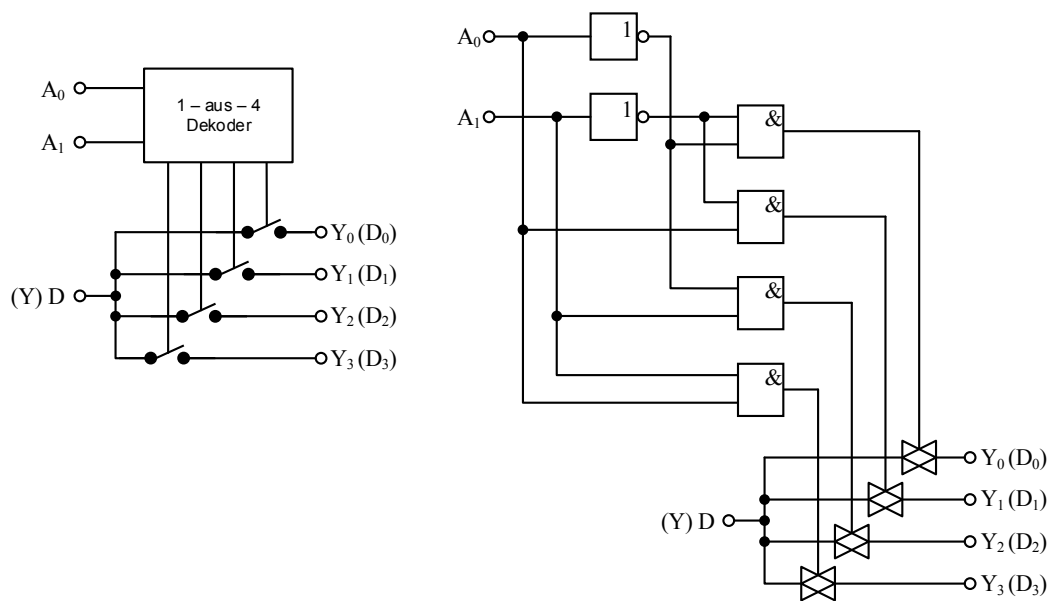


Bild 10.8: Analog-Multiplexer/Demultiplexer mit Transfer-Gattern.

die auf der linken Seite zu einem gemeinsamen Punkt verbunden sind. Die Bezeichnung “analog“ bezieht sich hierbei auf die Eigenschaft der Schalter, mit denen man (beliebige) analoge Spannungspegel übertragen kann. Als Multiplexer wird die Schaltung mit D_0 bis D_3 als Eingänge und Y als Ausgang betrieben. Wird jedoch Y zum Eingang D und die Eingänge D_0 bis D_3 als Ausgänge Y_0 bis Y_3 verwendet spricht man von einem Demultiplexer. In der schaltungstechnischen Realisierung in integrierten Schaltungen werden die eingezeichneten mechanischen Schalter durch Transfer-Gatter nach Bild 10.8b ersetzt.

11 Digital–Analog- und Analog–Digital–Wandler

Zwischen der analogen und digitalen Welt werden Schnittstellen zur Wandlung der Signale in die entsprechende analoge bzw. digitale Form benötigt. Diese Datenwandler sind sogenannte Analog- Digital (ADC–Analog-to-Digital Converter) bzw. Digital–Analog Wandler (DAC–Digital-to- Analog Converter). Aus der Alltagserfahrung sind digitale gespeicherte Signale z.B. auf Compact Discs (CD), im PC als MP3–Daten gut bekannt. Durch DAC's werden diese digitalen Daten in analoge Größen wie z.B. Spannungen oder Ströme umgesetzt. Besondere Bedeutung haben die Datenwandler durch den Einsatz von Computern in der Verfahrenstechnik zur Überwachung, Erfassung und Regelung von Prozessen im Industrie- und Heimbereich erlangt. Ein universelles Steuer- und Regelsystem besteht aus den Komponenten: Sensor mit Signalwandler, ADC, Computer, DAC und Stellglied. Die ADC und DAC erfassen bzw. erzeugen in der Regel Spannungen, die durch entsprechende Signalwandler an die entsprechende analoge Messgröße angepasst werden müssen. Sowohl DAC als auch ADC lassen sich grundsätzlich in drei Verfahren unterscheiden:

1. Parallelverfahren
2. Wägeverfahren
3. Zählverfahren

Parallelverfahren sind sehr genau, schnell und aufwendig, während Zählverfahren langsam und kostengünstig sind. Die Wägeverfahren werden für Wandler mit mittlerer Auflösung (8 – 16 Bit) und Schnelligkeit (0.5 – 10 μ s) eingesetzt.

11.1 Digital–Analog–Wandler

Lernziele

- Kennenlernen der grundsätzlichen Schaltungen zur Digital–Analog Wandlung
- Verstehen der unterschiedlichen Verfahren der Digital–Analog Wandlung
- Kennenlernen der typischen Fehlern bei Digital–Analog Wandlern

11.1.1 Grundlagen der Digital–Analog Wandlung

Digitale Werte lassen sich einfach durch Zuordnung von diskreten Werten in Form einer Tabelle darstellen. Als Beispiel betrachten wir Tabelle 11.1. Sie enthält für $2^4 = 16$

Binärzahl	z	$f(z)$
0000	0	1
0001	1	1,40673
0010	2	1,74313
0011	3	1,95105
0100	4	1,99453
0101	5	1,86606
0110	6	1,58785
0111	7	1,208
1000	8	0,79218
1001	9	0,4123
1010	10	0,13404
1011	11	0,00549
1100	12	0,0489
1101	13	0,25675
1110	14	0,59311
1111	15	0,99981

Tabelle 11.1: Digitale Darstellung der diskreten Werte z und der zugeordneten Werte der Funktion $f(z) = 1 + \sin(2\pi z/15)$.

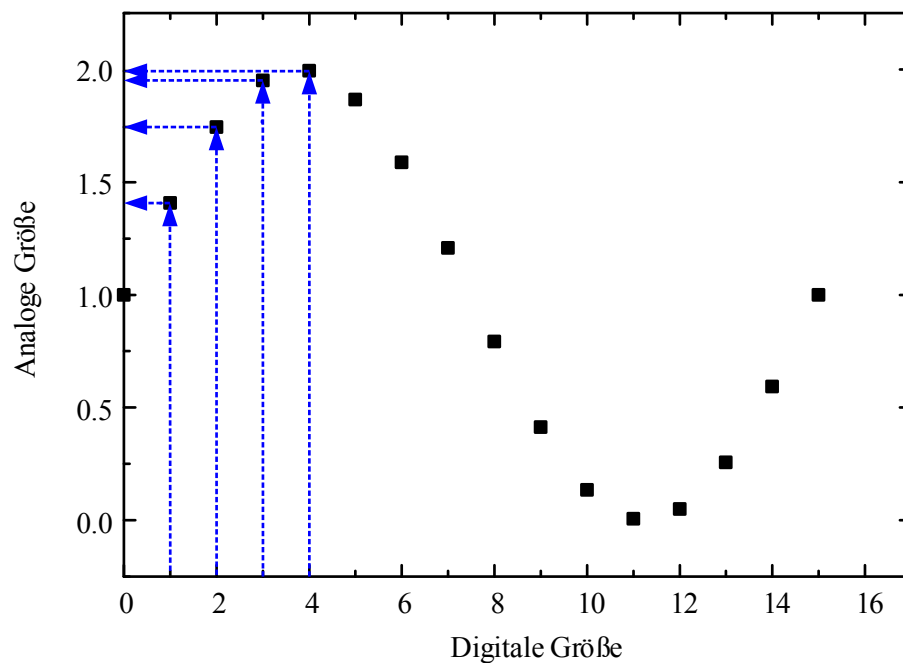


Bild 11.1: Zuordnung der digitalen Größe zu einer analogen Größe entsprechend Tabelle 11.1.

(4 Bit) diskrete Werte z die zugehörigen Werte der Funktion $f(z) = 1 + \sin(2z/15)$. In der Tabelle sind neben der dezimalen Darstellung der Werte z auch ihre binäre Darstellung gezeigt. Ordnen wir nun den Werten z und $f(z)$ bestimmte Segmente der x- und y-Achse einer Grafik zu, entsteht eine Folge von diskreten Werten (Bild 11.1), die man zu einer Kurve, d.h. einer analogen Information verbinden kann. Damit haben wir eine Digital–Analog Wandlung durchgeführt. Das entstandene Analogsignal ist ein diskret abgestuftes Signal, das sich aus der Anzahl der digitalen Abstufungen, also der Bitzahl ergibt. Damit lassen sich die wichtigsten Kenngrößen für DAC einführen. Diese Kenngrößen werden auch für die Analog–Digital Wandler benutzt, wobei sich die Abstufungen auf die entsprechende Wandlung der analogen Eingangswerte in digitale Werte beziehen. Die Anzahl der Abstufungen in Bit spezifiziert die Auflösung eines DAC. Somit hat ein DAC mit einer Auflösung von 4 Bit entsprechend $2^4 = 16$ Abstufungen oder ein DAC mit 10 Bit Auflösung hat $2^{10} = 1024$ Abstufungen. Die feinste analoge Abstufung, die dem LSB (Least Significant Bit $\equiv$ niederwertigstes Bit) entspricht, charakterisiert den Quantisierungsfehler ($\pm 0,5$ LSB) eines DAC. Die Wandlungsgeschwindigkeit gibt die Zahl der digital–analog Wandlungen pro Sekunde an und liegt je nach Wandlungsverfahren zwischen einigen Wandlungen/Sekunde und einigen 10^7 bis 10^8 Wandlungen/Sekunde. Häufig wird auch die Wandlungszeit angegeben. Im Folgenden wollen wir uns hier auf das Wägeverfahren konzentrieren, da sowohl das Parallelverfahren wie auch das Zählverfahren im aufgrund der erheblichen Nachteile heute praktisch keine Bedeutung mehr beim Einsatz von DACs haben.

11.1.2 Digital–Analog Wandlung nach dem Wägeverfahren

Beim Wägeverfahren wird jedem Bit ein gewichteter analoger Wert zugeordnet. Die Wichtung erfolgt entsprechend dem verwendeten Binärkode. Betrachten wir z.B. vier Bit dann können wir jeweils einem Bit n den binär gewichteten Wert (2^{n-1}) wie folgt zuordnen: 1 (1. Bit, d.h. LSB), 2 (2. Bit), 4 (3. Bit) und 8 (4. Bit, d.h. MSB – Most Significant Bit $\equiv$ höchstwertiges Bit). Entsprechend dem Binärkode können wir $2^4 = 16$ abgestufte analoge Werte den Binärzahlen 0000 bis 1111 zuordnen. So entspricht der Binärzahl 1001 ein analoger Wert von 9.

Bild 11.2 zeigt das prinzipielle Blockschaltbild eines DAC nach dem Wägeverfahren. Jedem Bit ist ein Schalter zugeordnet, der eine Referenzspannung auf ein Widerstandnetzwerk schaltet. Die Widerstände im Netzwerk sind entsprechend dem Code bewertet und somit werden Teilströme bzw. Teilspannungen erzeugt, die durch einen Verstärker summiert werden. D.h. einem digitalen Code am Eingang des DAC wird eine proportionale analoge Ausgangsspannung zugeordnet. Bild 11.3 zeigt den prinzipiellen Schaltplan eines DAC nach dem Wägeverfahren. Für einen 4–Bit DAC sind die Widerstände entsprechend der Wertigkeit $2^{n-1} \cdot R$ der Anzahl Bit n ($n = 4$) im Binärkode wie folgt bemessen: 1. Bit = R (MSB), 2. Bit = $2R$, 3. Bit = $4R$, 4. Bit = $8R$ (LSB). Der Widerstandswert R kann in weiten Grenzen frei gewählt werden und richtet sich nach der Größe von U_{ref} und der Auslegung der Schaltung. In der Regel beträgt er 1 k Ω bis 10 k Ω .

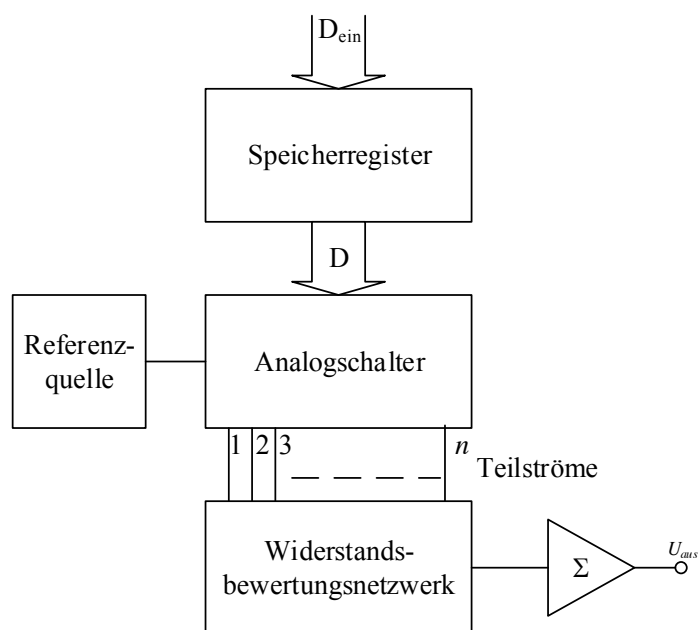


Bild 11.2: Blockschaltbild eines Digital–Analog Wandlers nach dem Wägeverfahren.

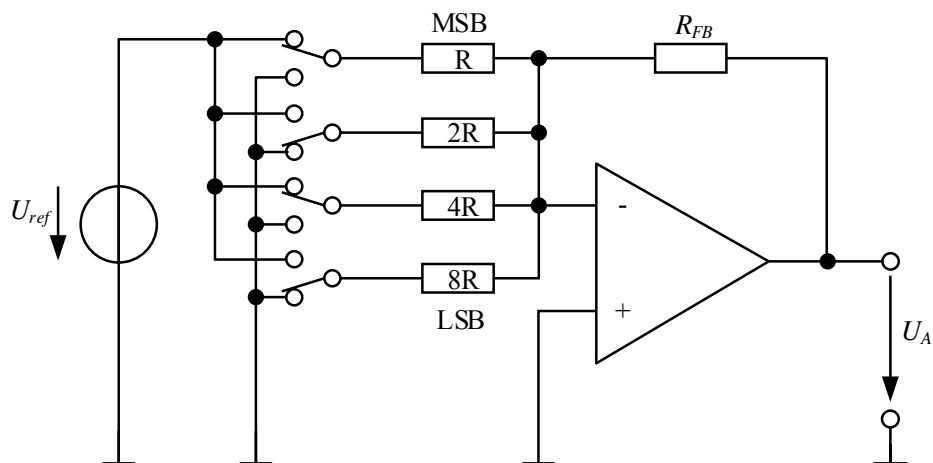


Bild 11.3: Prinzipschaltung eines 4–Bit DAC nach dem Wägeverfahren mit binär gewichtetem Widerstandsteiler, Referenzspannungsquelle und Summationsverstärker.

Im Folgenden sollen die Teilströme, die durch die Schalter eingestellt werden können berechnet werden. Die Schalterfunktion soll durch die Variable b_{n-1} beschrieben werden, wobei n dem entsprechenden Bit zugeordnet ist. Es fließt ein Strom in den Summenpunkt des Verstärkers (virtuelle Masse) bei $b_n = 1$. Bei $b_n = 0$ liegt der Widerstand an Masse. Damit lässt sich für den 4 Bit DAC der Strom in den Summenpunkt des Verstärkers in Form folgender Summe schreiben:

$$I_S = \frac{b_0 \cdot U_{ref}}{8 \cdot R} + \frac{b_1 \cdot U_{ref}}{4 \cdot R} + \frac{b_2 \cdot U_{ref}}{2 \cdot R} + \frac{b_3 \cdot U_{ref}}{R} \quad (11.1)$$

Mit einer Vereinfachung der Gleichung und einer Erweiterung des Nenners mit 2^{n-1} für das jeweilige Bit ergibt sich folgender allgemeine Ausdruck für den Strom:

$$I_S = \frac{U_{ref}}{2^{n-1} \cdot R} \cdot \sum_{i=0}^{n-1} 2^i \cdot b_i \quad (11.2)$$

Damit ergibt sich eine Ausgangsspannung von

$$U_{Ausgang} = -R_{FB} \cdot I_S \quad (11.3)$$

Die Ausgangsspannung des DAC ist somit proportional dem Binärkode, der die Widerstände auf die Referenzspannungsquelle schaltet. In dieser Diskussion wurden die Serienwiderstände der Schalter und der Innenwiderstand der Referenzspannungsquelle vernachlässigt. Ein Nachteil des diskutierten einfachen Widerstandteilers besteht darin, dass mit der Erhöhung der Auflösung um jedes weitere Bit der Widerstandswert des neuen Widerstandes doppelt so groß sein muss wie der größte Teilwiderstand. Somit ergeben sich sehr schnell große Widerstandswerte. Ein weiterer Nachteil ist die von der binären Eingangsinformation abhängige unterschiedliche Belastung der Referenzspannungsquelle.

Wägeverfahren mit R–2R–Netzwerken

Diese Nachteile lassen sich durch den Aufbau von sogenannten R–2R–Teilerketten umgehen, die heute in allen integrierten DAC's eingesetzt werden. Die R–2R–Teilerketten lassen sich einfach zu Netzwerken mit großer Genauigkeit (16 Bit) zusammenfügen und sind wesentlich preiswerter, da ihre Kalibrierung mit Laserabgleich wesentlich einfacher ist. Eine entsprechende Schaltung eines 8 bit DACs ist in Bild 11.4 dargestellt. Durch den Rückkopplungswiderstand R_{FB} (Feed–Back) liegen beide Eingänge des Operationsverstärkers, der als invertierender Verstärker betrieben wird, auf Massepotential. Damit ist die Belastung der Referenzspannungsquelle unabhängig von der Schalterstellung, d.h. vom Wert des binären Eingangssignals. Die Ströme addieren sich damit linear am Eingang des Operationsverstärkers. Die einzelnen Strompfade sind durch ein R–2R Leiternetzwerk binär gewichtet. Die in Bild 11.4 gezeigten Schalterstellungen entsprechen

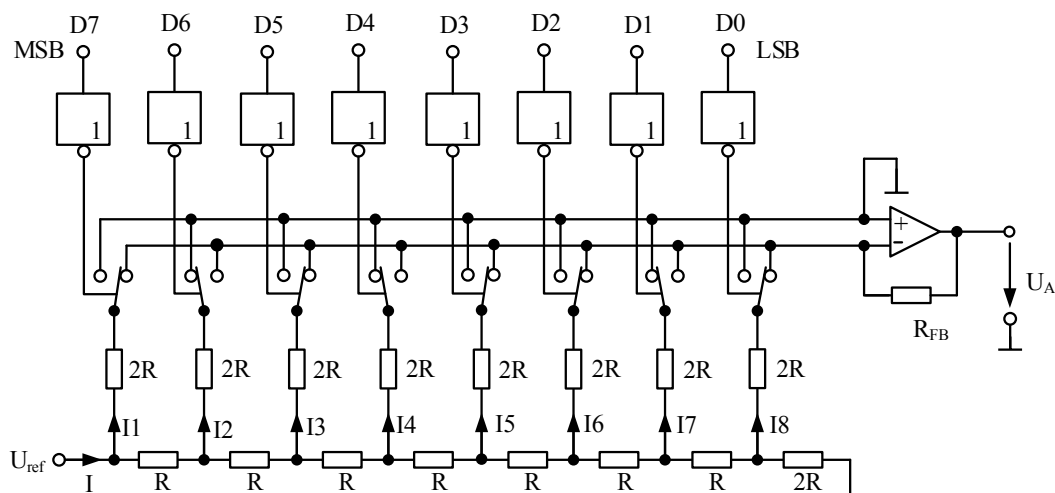


Bild 11.4: 8 Bit DAC mit R-2R-Netzwerk.

der binären Zahl “10101101“ $\hat{=}$ 173 dezimal. Die Ausgangsspannung U_A ist bei einer positiven Referenzspannung negativ:

$$U_A = -U_{ref} \cdot \frac{R_{FB}}{2^8 \cdot R} \cdot [1 \cdot D_0 + 2 \cdot D_1 + 4 \cdot D_2 + 8 \cdot D_3 + 16 \cdot D_4 + 32 \cdot D_5 + 64 \cdot D_6 + 128 \cdot D_7] \quad (11.4)$$

Sie ist unabhängig vom Absolutwert R , wenn auch der Rückkopplungswiderstand $R_{FB} = R$ ist. Die Werte von D_0 bis D_7 sind entsprechend der binären Information entweder 1 oder 0. Damit ergibt sich für die gezeichnete Schalterstellung in Bild 11.4:

$$U_A = -U_{ref} \cdot \frac{1}{256} \cdot [1 \cdot 1 + 2 \cdot 0 + 4 \cdot 1 + 8 \cdot 1 + 16 \cdot 0 + 32 \cdot 1 + 64 \cdot 0 + 128 \cdot 1] = -U_{ref} \cdot \frac{173}{256} \quad (11.5)$$

Dieses Wägeverfahren lässt sich besonders gut in der CMOS-Technik verwirklichen, da nur kleine Ströme notwendig sind und Transfer-Gatter als schnelle Schalter mit einem hohen Ein/Aus-Widerstandsverhältnis zur Verfügung stehen.

11.1.3 Genauigkeit von Digital-Analog-Wandlern

Man unterscheidet bei Digital-Analog-Wandlern grundsätzlich zwischen statischen und dynamischen Fehlern. Statische Fehler sind zum einen Nullpunktfehler, die durch Leckströme der Transistoren und Offsetfehler des Operationsverstärkers begründet sind, und zum andern Fehler beim Vollausschlag des Wandlers. Diese Fehler entstehen in erster Linie durch die endlichen Restwiderstände der Transistoren im eingeschalteten Zustand. Einen zusätzlichen Einfluss hat dann noch die Genauigkeit des Rückkopplungswiderstandes, der ja die Verstärkung bestimmt. Diese Fehler lassen sich durch einen Abgleich der

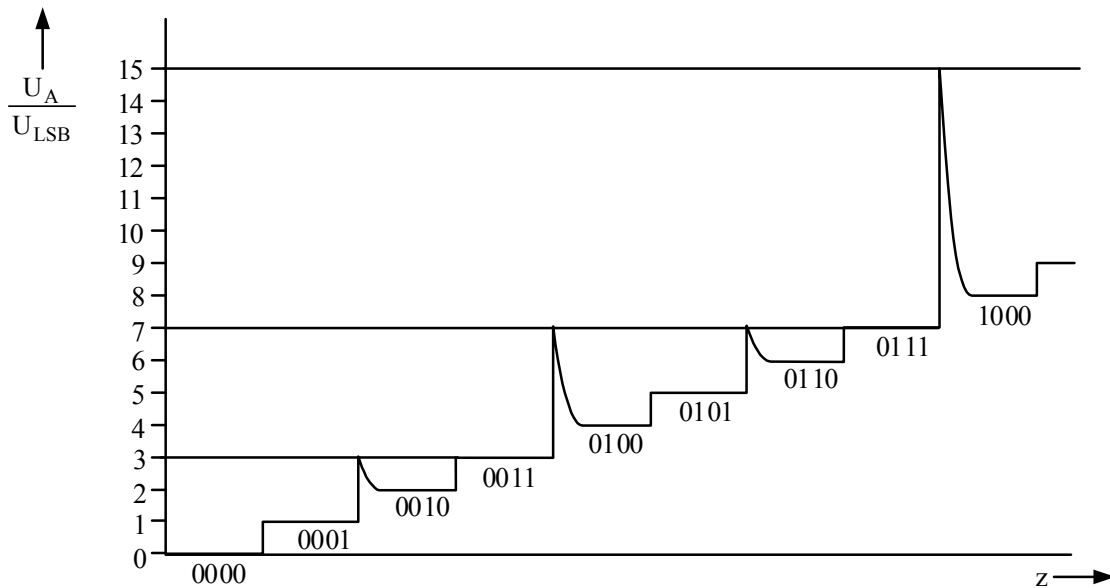


Bild 11.5: Darstellung der 'Glitches' am Ausgang eines DACs.

beiden Arbeitspunkte über Offseteinstellung und Korrektur der Verstärkung weitgehend beseitigen. Nicht abgleichbar sind dagegen Fehler in der Linearität der analogen Ausgangsgröße. Diese können durch schlecht abgeglichene Widerstände R entstehen. Der Fehler in den einzelnen Zweigen sollte insgesamt nicht größer sein, als der Strom durch den Zweig des niederwertigsten bits des Wandlers, nämlich $\pm \frac{1}{2}$ LSB. Dynamische Fehler bei DACs sind kurze Störimpulse (sogenannte Glitches) am analogen Ausgang. Die Ursache dieser Impulse liegt am nicht gleichzeitigen Umschalten der Schalter. Mögliche Gründe hierfür sind:

- nicht gleichzeitig eintreffende Eingangssignale an den digitalen Eingängen,
- Unterschiede in den Umschaltzeiten der Transistoren zwischen Ein- und Ausschalten.

Die Beseitigung dieser Störimpulse ist nicht einfach und wird bei langsameren D/A-Wandlern bereits auf den Chips so weit wie möglich vorgenommen. Der erste Einfluss für die Störungen kann durch den Einsatz von zusätzlichen Eingangsregistern auf den DAC erheblich reduziert werden, während die unterschiedlichen Umschaltzeiten der Transistoren praktisch nicht umgangen werden können. Aus Bild 11.5 erkennt man, dass Glitches durch die Betätigung von mehr als einem Schalter zur Einstellung der nächst höheren binären Zahl auftreten und mit der Anzahl der Stellen der Binärzahl zunehmen. In vielen Datenblättern ist die maximale Energie der Glitches angegeben.

11.2 Analog–Digital Wandlung

Lernziele

- Kennenlernen der Funktion von Schaltungen zur Analog–Digital Wandlung
- Verstehen der unterschiedlichen Verfahren der Analog–Digital Wandlung
- Kennenlernen der Genauigkeit und der typischen Fehler bei Analog–Digital Wandlern

11.2.1 Grundlagen der Analog–Digital Wandlung

Wie bereits besprochen, sollen die verschiedensten sich in der Zeit ändernden Analogsignale bereits durch Sensoren und entsprechende Wandler in ein zeitlich veränderliches analoges Spannungssignal umgewandelt worden sein. Hier soll nur die Wandlung der entsprechenden analogen Spannungssignale in digitale Werte besprochen werden. Da sich die analogen Spannungssignale sowohl in der Zeit als auch in ihrer Amplitude ändern, müssen sie sowohl in der Zeit als auch in der Amplitude diskretisiert werden. Der prinzipielle Prozess der zeitlichen Diskretisierung ist in Bild 11.6 dargestellt. Das Spannungssignal wird zu den Zeiten $t_1, t_2, \dots, t_n$ abgetastet. Dieser Vorgang wird als “Abtasten” bzw. “Sampling” bezeichnet. Nachdem die Spannungssegmente (auch als “samples” bezeichnet) zu den entsprechenden Zeitpunkten erzeugt wurden, erfolgt die Umsetzung der Amplitudenwerte in die Digitalwerte.

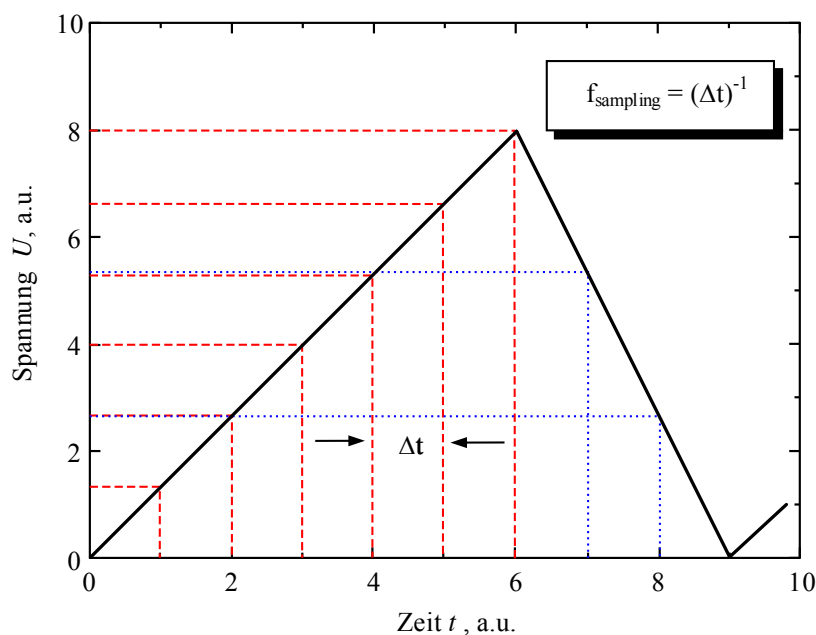


Bild 11.6: Darstellung der 'Glitches' am Ausgang eines DACs.

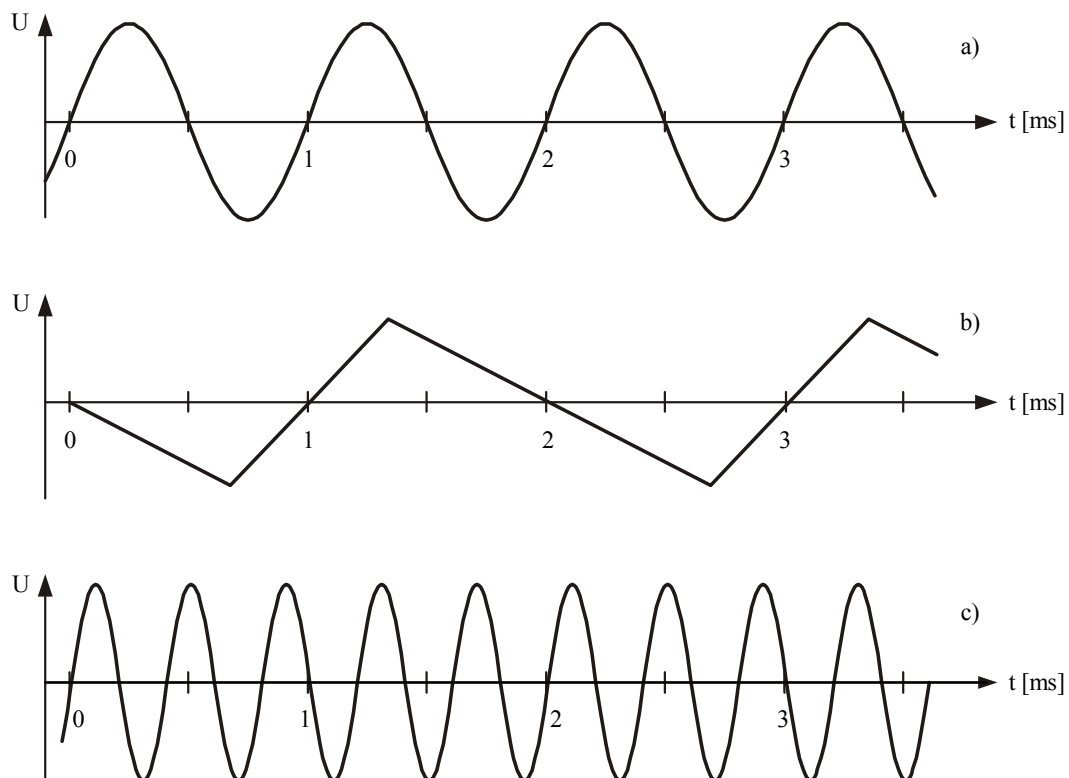


Bild 11.7: Abtastung zweier Signal (1 kHz: Kurve a und 2,5 kHz: Kurve c mit einer Abtastfrequenz von 1,5 kHz. Kurve b zeigt das rekonstruierte Signal, das aus den abgetasteten Spannungswerten zusammengesetzt wurde).

Die entscheidende Frage ist nun, wie groß die Abtastfrequenz mindestens sein muss, damit das ursprüngliche analoge Signal wieder rekonstruiert werden kann.

Abtast-Theorem (Nyquist-Theorem)

Zur Veranschaulichung sind in Bild 11.7 zwei Sinussignale mit den Frequenzen von 1 kHz (Kurve a) und 2,5 kHz (Kurve c) dargestellt, die alle 0,67 ms (also mit einer Frequenz von 1,5 kHz) abgetastet werden. Die jeweiligen zugehörigen analogen Spannungswerte sind in Kurve b als Punkte eingezeichnet und mit Geraden verbunden wurden. In Kurve b ist die zeitliche Information (bzw. die Frequenzinformation) der beiden Ausgangssignale verloren gegangen. Ein weiteres Beispiel ist in Bild 11.8 dargestellt. Das Ausgangssignal mit einer Frequenz von 1 kHz wird mit verschiedenen Frequenzen abgetastet. Bei einer Abtastfrequenz von 10 kHz lässt sich das Signal gut rekonstruieren, während mit abnehmender Abtastfrequenz die Ergebnisse immer schlechter werden. Bei einer Abtastfrequenz von 625 Hz ergibt sich bei der Rekonstruktion ein völlig falscher zeitlicher Signalverlauf. Die minimal erforderliche Abtastfrequenz wurde 1929 von Nyquist berechnet und wird in folgendem Theorem definiert:

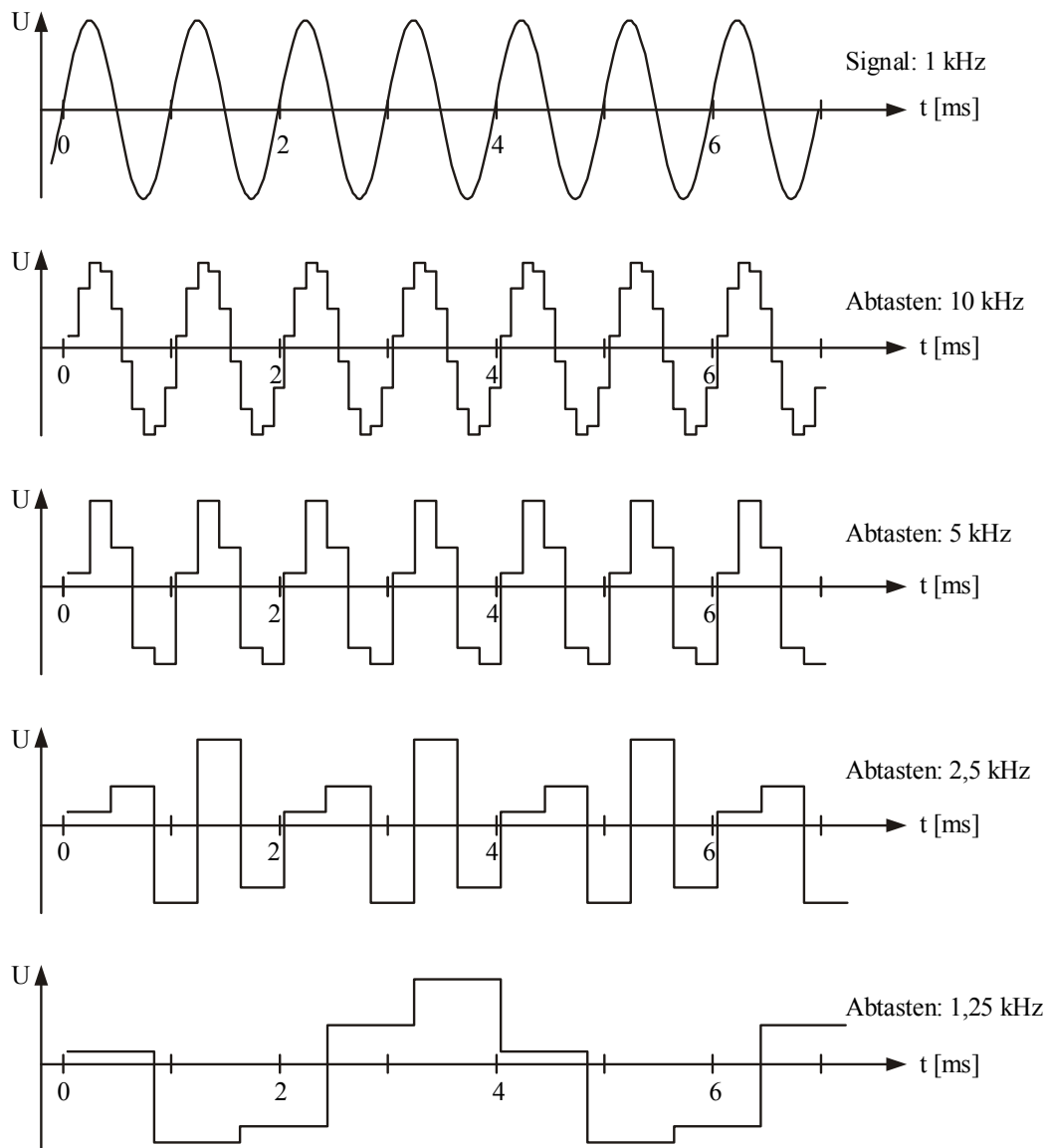


Bild 11.8: Zeitliche Abtastung eines 1 kHz Signals mit verschiedenen Abtastfrequenzen.

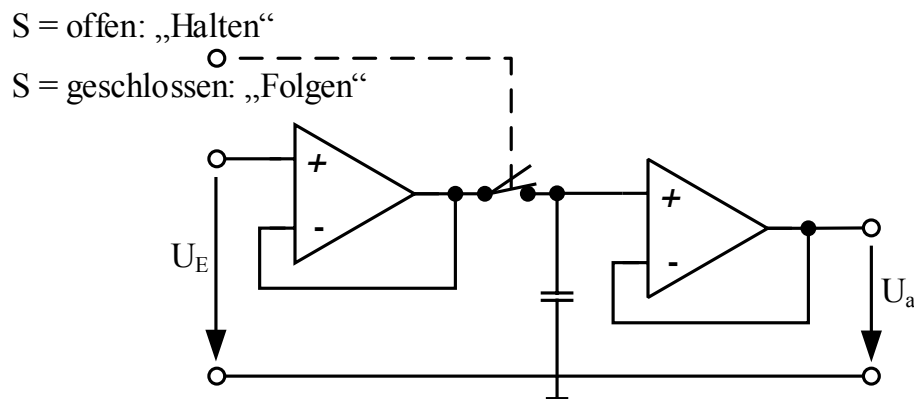


Bild 11.9: Zeitliche Abtastung eines 1 kHz-Signals mit verschiedenen Abtastfrequenzen.

Der zeitliche Verlauf eines Signals mit den Frequenzkomponenten f_{max} kann durch Abtastung mit einer Frequenz von mindestens $2 \cdot f_{max}$ eindeutig wiedergegeben werden.

Für das Beispiel in Bild 11.8 muss also die Abtastfrequenz mindestens 2 kHz betragen. Die Auswahl der richtigen Abtastfrequenz und der zugehörigen Filter im Signalzweig bestimmt wesentlich die zeitliche Diskretisierung der Spannungssignale und muss deshalb mit großer Sorgfalt erfolgen. Da die Umwandlung des analogen Eingangssignals in ein digitales Ausgangssignal eine mehr oder minder lange Zeit benötigt, ist es für die Genauigkeit der Wandlung erforderlich, das zeitlich veränderliche Eingangssignal zu einem bestimmten Zeitpunkt abzutasten, und den analogen Wert für die Dauer der Wandlung konstant zu halten. hierzu werden sogenannte Abtast- und Halte-Schaltungen nach Bild 11.9 eingesetzt. Die Schaltung besteht aus zwei hintereinander geschalteten Operationsverstärkern. Im Modus „Folgen“ ist der Schalter geschlossen und der Kondensator zwischen den beiden Operationsverstärkern wird auf den Wert des Eingangssignals aufgeladen. Während des Haltevorgangs ist der Ausgang des ersten Operationsverstärkers vom Eingang des zweiten abgetrennt. Um nun die Eingangsspannung des zweiten Verstärkers für die A/D-Wandlung konstant zu halten, sollte der Kondensator sich möglichst wenig entladen. Daher muss sowohl der Eingangsstrom des verwendeten Operationsverstärkers wie auch der Leckstrom des Kondensators besonders klein gehalten werden. Der Spannungsverlust ΔU während der Wandelzeit muss kleiner als $\frac{1}{2} \cdot U_{LSB}$ sein, damit kein Fehler bei der Wandlung auftritt. Bei einem 12 bit Wandler entspricht dies $\Delta U \leq 0,5 \cdot U_{Emax}/4095$.

11.2.2 Parallel-Wandler (Flash)

Das Parallelverfahren ist das schnellste, zugleich aber auch das schaltungstechnisch aufwendigste. Hierbei werden alle n bits der digitalen Ausgangsinformation gleichzeitig erzeugt. Man bezeichnet diese Art von A/D-Wandler deshalb auch als „Flash“-Wandler. Um dies zu erreichen, benötigt man auf der analogen Eingangsseite einen Spannungs-

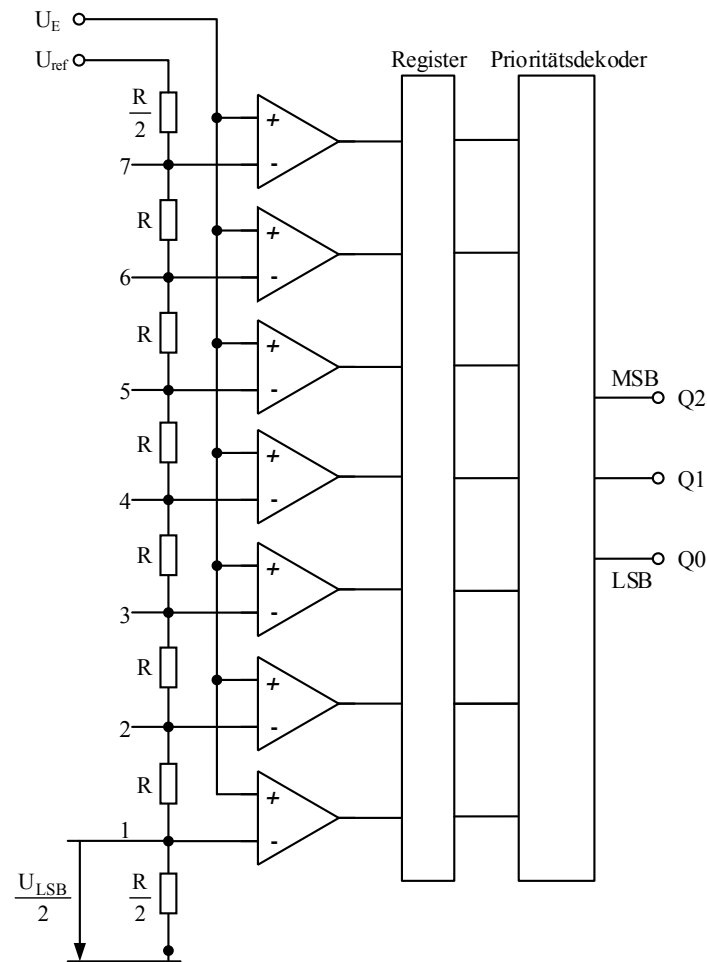


Bild 11.10: Blockschaltbild eines 3 bit Analog–Digital–Wandlers nach dem Parallelverfahren.

teiler mit $(2^n - 1)$ exakt gleichen Stufen und ebenso viele Komparatoren. In Bild 11.10 ist ein 3 bit Parallel – A/D–Wandler dargestellt. Die Abgriffe an den Punkten 1 bis 7 des Spannungsteilers bilden die Spannungsreferenzen für die sieben Komparatoren. Die niedrigste Referenzspannung entspricht der Hälfte der Quantisierungsstufe der niederwertigsten bits, nämlich $\frac{1}{2} \cdot U_{LSB}$. Die nächst größere Referenzspannung ist um $1 \cdot U_{LSB}$ höher usw.. Abhängig von der Größe des analogen Eingangssignals schalten alle Komparatoren, deren Referenzspannung kleiner oder gleich der Eingangsspannung ist, ein. Die Ausgangsinformation aller Komparatoren wird dann zu einem genau vorgegebenen Zeitpunkt in einem digitalen Register zwischengespeichert und durch einen Prioritätsdecoder (wie z.B. Bild 10.3) in die gewünschte binäre Form gebracht. Da sich bei diesem Verfahren die Anzahl der benötigten Komparatoren mit jedem zusätzlichen bit der Auflösung verdoppelt (für 8 bit $2^8 - 1 = 255$ Komparatoren und für 9 bit $2^9 - 1 = 511$ Komparatoren), wird der schaltungstechnische Aufwand sehr schnell außerordentlich groß. Mit dem Parallelverfahren werden mit Standardtechnologien wie heute Abtast-

Bit	Test	Vergleich $U_{max} = 64$	(für Setzen	Binärzahl
5 (MSB)	$\left(\frac{1}{2}\right) \cdot U_{max}$	$x \geq 32$	Ja	1 ? ? ? ? ?
4	$\left(\frac{1}{2} + \frac{1}{4}\right) \cdot U_{max}$	$x \geq 32 + 16$	Nein	1 0 ? ? ? ?
3	$\left(\frac{1}{2} + 0 + \frac{1}{8}\right) \cdot U_{max}$	$x \geq 32 + 8$	Ja	1 0 1 ? ? ?
2	$\left(\frac{1}{2} + 0 + \frac{1}{8} + \frac{1}{16}\right) \cdot U_{max}$	$x \geq 32 + 8 + 4$	Nein	1 0 1 0 ? ?
1	$\left(\frac{1}{2} + 0 + \frac{1}{8} + 0 + \frac{1}{32}\right) \cdot U_{max}$	$x \geq 32 + 8 + 2$	Ja	1 0 1 0 1 ?
0	$\left(\frac{1}{2} + 0 + \frac{1}{8} + 0 + \frac{1}{32} + \frac{1}{64}\right) \cdot U_{max}$	$x \geq 32 + 8 + 2 + 1$	Ja	1 0 1 0 1 1

Tabelle 11.2: Bitweise Annäherung des Vergleichssignals für einen 6 Bit ADC nach dem Wägeverfahren.

frequenzen von $f_A = 3$ GHz, d.h. Wandlungsraten von 3000 MSPS (MegaSamples Per Second) bei Wandlern mit einer Auflösung von 8 bit erreicht. Die Wandlungszeit (conversion time) für $f_A = 3$ GHz ist also $\tau_A = 0,33$ ns. Für sehr schnelle Oszilloskope werden für die Eingangssignale Abtastfrequenzen von bis zu 10 GHz angegeben. Über die eingesetzten DACs erhält man aber vorerst keine tiefer gehende Information vom Hersteller.

11.2.3 Analog-Digital Wandlung mit sukzessiver Approximation

Hier soll die Analog–Digital Wandlung (ADC) nach dem Wägeverfahren besprochen werden. Das Wägeverfahren lässt sich leicht am Beispiel einer Balkenwaage veranschaulichen. Die unbekannte Größe (entspricht der analogen Eingangsspannung) wird auf die Waage aufgelegt und durch genau bemessene Gewichte wird versucht, die Waage ins Gleichgewicht zu bringen. Betrachten wir ein Beispiel eines ADC mit der Ausgangsspannung U_{max} und einer Wichtung mit 6 diskreten Werten entsprechend dem Binärcode. Damit ist der geringste Wert (LSB) auf $1/2^6 = 1/64$ festgelegt. Die unbekannte Eingangsspannung soll im Beispiel einem Wert von 43 (für $U_{max}=64$) haben. Tabelle 11.2 zeigt die einzelnen Schritte des Wägeverfahrens. Den Verlauf der Vergleichsgröße (z.B. der Vergleichsspannung) im Verlauf des Verfahrens zeigt Bild 11.11. Im ersten Schritt wird das höchstwertige Bit also die Hälfte des Maximalwertes U_{max} angelegt und verglichen. Ist die Eingangsspannung U_E größer, wird das bit gesetzt (d.h. bei der Waage bleibt das Gewicht liegen). In jedem darauffolgenden Schritt wird der zuletzt benutzte Wert halbiert. Wenn dabei die Summe der Vergleichswerte den Wert der Eingangsgröße überschreitet, wird das bit des letzten Teilwerts nicht gesetzt (d.h. bleibt nicht auf der Waage liegen). Das ist in Tabelle 11.2 bei bit 2 und bit 4 der Fall. Bild 11.11 zeigt die gleichmäßige Annäherung des Ausgangswertes an den Wert der Eingangsgröße. Deshalb wird dieses Verfahren auch als Verfahren der sukzessiven Approximation bezeichnet. Die Annäherung der gewichteten Vergleichswerte an den Eingangswert kann nur mit

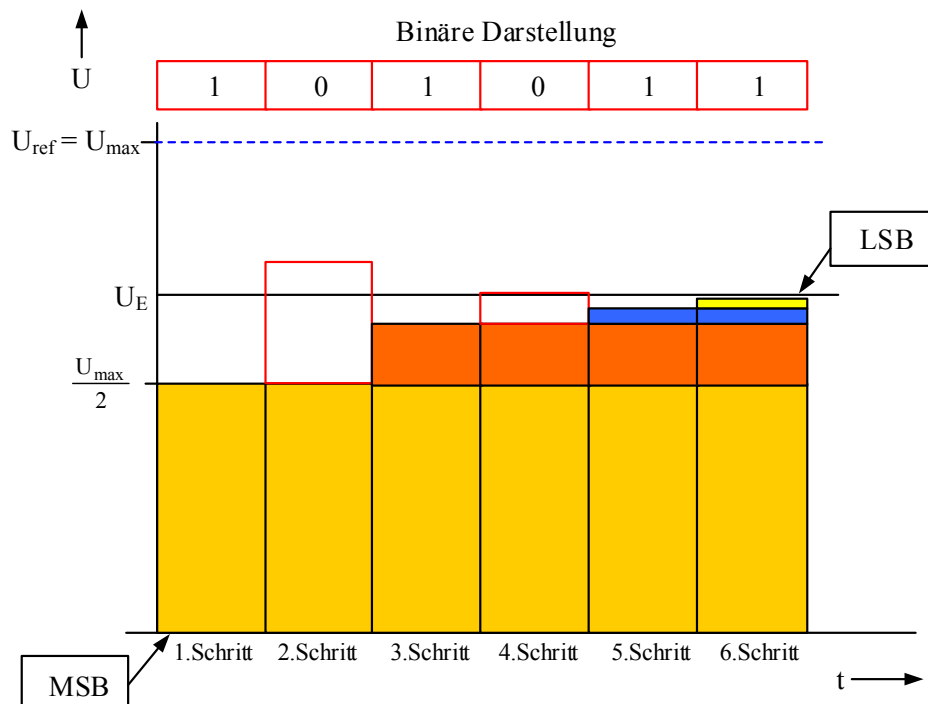


Bild 11.11: Zeitliche Abfolge der sukzessiven Annäherung der Vergleichsspannung (d.h. der Ausgangsspannung des DAC) in einem 6 bit ADC.

einer Genauigkeit von $\pm \frac{1}{2} \cdot \text{LSB}$ erfolgen. Das entspricht dem Quantisierungsfehler eines ADC. Die Auflösung eines ADC wird durch die Bitzahl der Abstufungen ausgedrückt in 2^n charakterisiert. Die erforderliche Gesamtzeit der Analog-Digital Wandlung wird als Wandlungszeit bezeichnet. Ausgehend von diesem Beispiel lässt sich leicht der prinzipielle Aufbau eines DAC nach dem Verfahren der sukzessiven Approximation ableiten.

Bild 11.12 zeigt die wesentlichen Komponenten. Der Komparator vergleicht die Eingangsspannung mit der gewichteten Vergleichsspannung. Die Vergleichsspannung kommt von einem Digital-Analog Wandler (DAC), der ebenfalls nach dem Wägeverfahren arbeitet. Ist die Eingangsspannung größer als die Vergleichsspannung so ist der Komparatorausgang logisch 1 und setzt das MSB auf 1. Der Vorgang des Setzens bzw. Löschs der Bits für den DAC wird im Sukzessiven Approximationsregister (SAR) durchgeführt. Das SAR besteht aus einem Schieberegister, das vom Ausgangssignal des Komparators gesteuert wird. Die digitalen Daten stehen am Ausgang des SAR als Binärzahlen im parallelen Format bereit. Der Taktgenerator stellt den Takt für die Wandlung bereit. Die Referenzspannungsquelle versorgt das Widerstandsnetzwerk des DAC. Da sich in der Regel die Eingangsspannung während der Umwandlungszeit ändert, benötigt man ein Abtast-Halte-Glied. Das S&H-Glied diskretisiert die Eingangsspannung im Zeitbereich und hält die Spannung am Komparatoreingang für die Dauer der Wandlungszeit konstant.

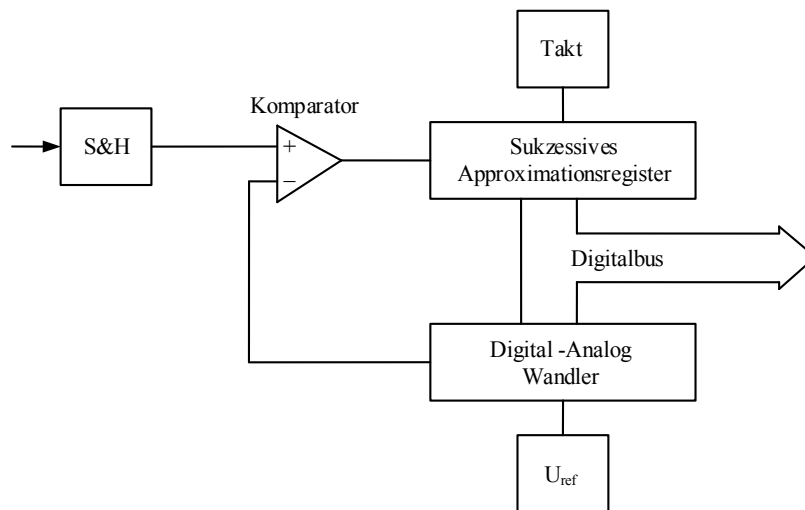


Bild 11.12: Blockschaltbild eines Analog-Digital Wandlers nach dem Verfahren der sukzessiven Approximation.

11.2.4 Analog-Digital Wandlung mit Integrationsverfahren

Bei den Integrationsverfahren unterscheidet man grundsätzlich zwei Varianten, die digitale Integration mit einem Zähler und die analoge Integration mit einem analogen Integrator. In dieser Vorlesung soll an dieser Stelle nur auf das einfachste analoge Integrationsverfahren, das Ein-Rampen-Verfahren (Single Slope), das in Bild 11.13 skizziert ist, eingegangen werden. In diesem Wandler ist ein Sägezahngenerator integriert, der von einer Referenzspannungsquelle angesteuert wird und eine Rampe mit konstanter Steigung liefert. Diese Sägezahnspannung wird als Referenzspannung für zwei analoge Komparatoren verwendet. Komparator 1 vergleicht die Eingangsspannung mit der Sägezahnspannung und Komparator 2 überprüft die Polarität der Sägezahnspannung, die von $-U_{ref}/2$ bis $+U_{ref}/2$ läuft. Die Ausgänge der beiden Komparatoren werden auf eine Äquivalenz-Schaltung geführt. Die Dauer D des Impulses am Ausgang G der Äquivalenzschaltung ist der zu messenden Spannung proportional. Den digitalen Ausgangswert erhält man durch das Zählen der Anzahl der Taktimpulse Z während der Dauer D . Als Takt verwendet man üblicherweise eine Rechteckschwingung, deren Frequenz über einen Schwingquarz stabilisiert wurde. Damit wird das Ergebnis der Wandlung zu:

$$Z = \frac{D}{T} = \tau \cdot f \cdot \frac{U_E}{U_{ref}} \quad (11.6)$$

Nach jeder Messung muss das Ergebnis in ein Ausgangsregister abgespeichert und der Zähler auf 0 zurückgesetzt werden. Die Erzeugung der Sägezahnschwingung geschieht durch einen Integrator, dessen Zeitkonstante durch ein RC-Glied bestimmt ist. Deshalb gehen Temperatur- und Langzeitdrift der Bauelemente voll in die Zeitkonstante und damit in die Messgenauigkeit ein. Eine Verbesserung der Genauigkeit erreicht man durch das Zwei-Rampen-Verfahren oder ein Zwei-Rampen-Verfahren mit automatischem Nullpunkt-Abgleich das auch als Quad-Slope-Verfahren bezeichnet wird. Diese

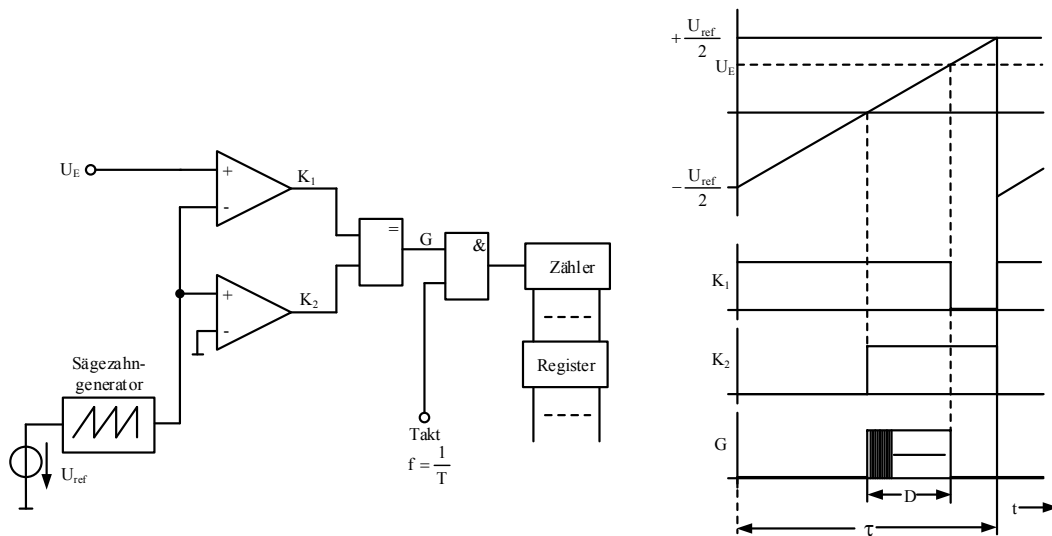


Bild 11.13: A/D-Wandler nach dem Ein-Rampen-Verfahren: a) Blockschaltbild und b) Signalverläufe im Wandler.

beiden und noch weitere Verfahren der Analog/Digital- Wandlung werden ausführlich in der Vorlesung “Integrierte Systeme und Schaltungen“ behandelt.

11.2.5 Genauigkeit von A/D-Wandlern

Abschließend sollen noch einige Betrachtungen zur Genauigkeit von A/D-Wandlern angestellt werden. Man unterscheidet grundsätzlich zwischen zwei Arten von Fehlern. Dies sind: 1. statische und 2. dynamische Fehler.

Statische Fehler

Ein statischer Fehler tritt systematisch durch die nur endliche Anzahl von auflösbaren Quantisierungsstufen auf. Er beträgt $\pm \frac{1}{2} \cdot \text{LSB}$. Daneben treten noch weitere, schaltungstechnisch bedingte Fehler auf. Die im Bild 11.14 eingezeichnete Gerade hat die Steigung 1 und geht exakt durch den Nullpunkt. Bedingt durch Offset-Fehler ist diese Kurve in Wirklichkeit jedoch nach rechts oder nach links verschoben und Verstärkungsfehler sorgen für eine mittlere Steigung, die von 1 abweicht. Diese Fehler lassen sich durch einen Abgleich im Nullpunkt durch Offsetkorrekturen und bei der maximalen Eingangsspannung $U_{E_{max}}$ durch Korrektur des Verstärkungsfaktors jedoch verkleinern. Nach diesem Abgleich bleiben nichtlineare Fehler, welche außer den unvermeidlichen Quantisierungsfehlern auftreten.

Totale Nichtlinearität

Die maximale Abweichung von der Geraden abzüglich des Quantisierungsfehlers bezeichnet man als totale Nichtlinearität (TN). Sie kann positiv oder negativ sein. Im

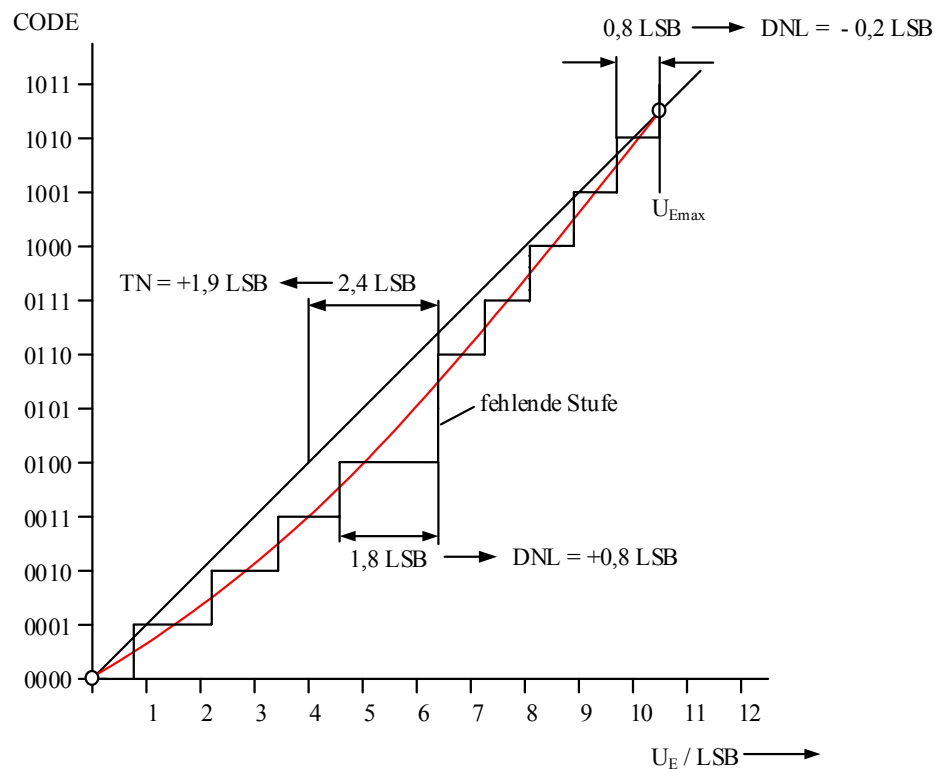


Bild 11.14: Darstellung verschiedener Abweichungen von der Linearität der Wandlung.

Beispiel in Bild 11.14 ist der Fehler bei der binären Eingangskombination 0 1 0 0 insgesamt 2,4 LSB. Subtrahiert man davon den Quantisierungsfehler 0,5 LSB erhält man als totale Nichtlinearität an dieser Stelle 1,9 LSB.

Differentielle Nichtlinearität

Die differentielle Nichtlinearität (DNL) gibt an, um welchen Wert die Breite der einzelnen Stufen vom Sollwert U_{LSB} abweicht. Beträgt die differentielle Nichtlinearität einer Stufe $DNL = 1 \cdot U_{\text{LSB}}$, bedeutet dies, dass der entsprechende binäre Zahlenwert des Ausgangssignals übersprungen wird. Man bezeichnet diesen Fall auch als "Missing Code".

Dynamische Fehler

Zusätzlich zu den statischen Fehlern treten dynamische Fehler während des Messvorganges auf, deren Ursache hier näher betrachtet werden soll. Um während der Wandlungszeit des A/D-Wandlers ein konstantes Eingangssignal zu erhalten, ist es notwendig, ein Abtast- und Halte-Glied vorzuschalten. Deshalb muss man zur Beurteilung der Genauigkeit die Eigenschaften des A/D-Wandlers und des Abtast- und Halte-Glieds zusammen betrachten. Das Abtast- und Halte-Glied muss wie bereits erwähnt, in der Lage sein, innerhalb der vorgegebenen Wandlungszeit des nachgeschalteten A/D Wand-

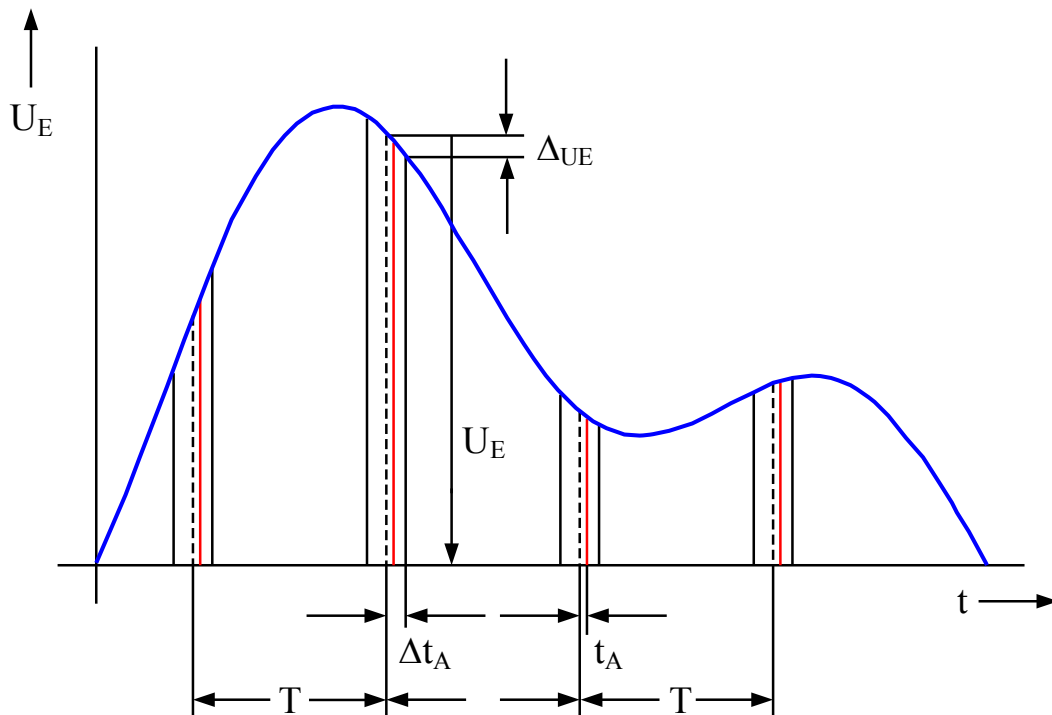


Bild 11.15: Darstellung der dynamischen Fehler bei A/D-Wandlern.

lers am Ausgang einen Spannungswert einzustellen, der auf $\pm \frac{1}{2} \cdot U_{LSB}$ genau mit der Eingangsspannung des Abtast- und Halte-Glieds übereinstimmt. Während der Haltezeit, d.h. während der Wandelzeit, muss die Eingangsspannung des A/D-Wandlers möglichst konstant gehalten werden. Die Ausgangsspannung des Abtast- und Halte-Glieds darf sich maximal um $-\frac{1}{2} \cdot U_{LSB}$ ändern, damit keine Fehler bei der Wandlung entstehen. Ein weiteres Problem sind Schwankungen des Abtastzeitpunktes wenn von "Folgen" auf "Halten" umgeschaltet bzw. die Wandlung im A/D-Wandler angestoßen wird. Diese Schwankungen werden als Apertur-Jitter Δt_A bezeichnet. Aufgrund einer Aperturzeit t_A kann es vorkommen, dass das Eingangssignal nicht genau zum Zeitpunkt t sondern um t_A verspätet abgetastet wird. Solange t_A konstant ist, hat dies keinen Einfluss auf das Ergebnis der Messung, da die Abstände der Abtastzeitpunkte konstant bleiben. Schwankt die Aperturzeit t_A jedoch um einen Wert Δt_A , so kann dies einen Fehler ΔU_E zur Folge haben, wie dies in Bild 11.15 gezeigt ist. Deshalb muss beim Entwurf von A/D-Wandlern darauf geachtet werden, dass der Apertur-Jitter Δt_A so kurz ist, dass ΔU_E immer kleiner $\frac{1}{2} \cdot U_{LSB}$ bleibt.

12 Literatur

Tietze/Schenk, Halbleiter-Schaltungstechnik, Springer Verlag 1999

Beuth, Elektronik 2, Bauelemente, Vogel Buchverlag 1998

Beuth, Elektronik 4, Digitaltechnik, Vogel Buchverlag 1998

Seifart/Beikirch, Digitale Schaltungen, Verlag Technik Berlin 1998

Hartl/Krasser/Pribyl/Söser/Winkler, Elektronische Schaltungstechnik, Pearson Studium 2008

Siegl, Schaltungstechnik, Springer Verlag 2004

Formelsammlung zur Klausur "Elektronische Schaltungen"

E 24 – Reihe: | 1,0 | 1,1 | 1,2 | 1,3 | 1,5 | 1,6 | 1,8 | 2,0 | 2,2 | 2,4 | 2,7 | 3,0 | 3,3 | 3,6 | 3,9 | 4,3
| 4,7 | 5,1 | 5,6 | 6,2 | 6,8 | 7,5 | 8,2 | 9,1 |

Diode:

$$I = I_S \cdot \left(e^{\frac{U}{n \cdot U_T}} - 1 \right) \quad \text{mit} \quad U_T = \frac{k_B \cdot T}{e}$$

Spannungsverstärkung und Ausgangswiderstand gelten bei Leerlauf am Ausgang

Bipolarer Transistor:

$$\text{Steilheit : } S = \frac{\partial I_C}{\partial U_{BE}} = \frac{I_{C,A}}{U_T}$$

$$\text{Kleinsignal-Eingangs-Widerstand: } r_{BE} = \frac{\beta}{S}$$

a) Emitterschaltung

$$\text{Kleinsignal-Spannungsverstärkung: } A = -S \cdot R_C$$

$$\text{Kleinsignal-Eingangswiderstand: } r_e = r_{BE}$$

$$\text{Kleinsignal-Ausgangswiderstand: } r_a = \left. \frac{\partial u_a}{\partial i_a} \right|_A = R_C$$

b) Emitterschaltung mit Stromgegenkopplung

$$\text{Kleinsignal-Spannungsverstärkung: } A = \left. \frac{u_a}{u_e} \right|_{i_a=0} \approx -\frac{SR_C}{1 + SR_E} \stackrel{SR_E \gg 1}{\approx} -\frac{R_C}{R_E}$$

$$\text{Kleinsignal-Eingangswiderstand: } r_e = \frac{u_e}{i_e} \approx r_{BE} + \beta R_E = r_{BE}(1 + SR_E)$$

$$\text{Kleinsignal-Ausgangswiderstand: } r_a = \frac{u_a}{i_a} \approx R_C$$

c) Kollektorschaltung

$$\text{Kleinsignal-Spannungsverstärkung: } A = \left. \frac{u_a}{u_e} \right|_{i_a=0} \approx \frac{SR_E}{1 + SR_E} \stackrel{SR_E \gg 1}{\approx} 1$$

$$\text{Kleinsignal-Eingangswiderstand: } r_e = \left. \frac{u_e}{i_e} \right|_{i_a=0} \approx r_{BE} + \beta \cdot R_E \stackrel{SR_E \gg 1}{\approx} \beta \cdot R_E$$

$$\text{Kleinsignal-Ausgangswiderstand: } r_a = \frac{u_a}{i_a} \approx R_E \parallel \left(\frac{R_g}{\beta} + \frac{1}{S} \right)$$

d) Basisschaltung

Kleinsignal-Spannungsverstärkung: $A = \frac{u_a}{u_e} \bigg|_{i_a=0} \approx \frac{\beta \cdot R_C}{r_{BE} + R_{BV}} \stackrel{r_{BE} \gg R_{BV}}{\approx} SR_C$

Kleinsignal-Eingangswiderstand: $r_e = \frac{u_e}{i_e} \approx R_E \parallel \left(\frac{1}{S} + \frac{R_{BV}}{\beta} \right) \stackrel{r_{BE} \gg R_{BV}}{\approx} R_E \parallel \frac{1}{S} \stackrel{R_E \rightarrow \infty}{\approx} \frac{1}{S}$

Kleinsignal-Ausgangswiderstand: $r_a = \frac{u_a}{i_a} \approx R_C$

Sperrschicht-Feldeffekttransistoren, selbstleitende Isolierschicht- Feldeffekttransistoren:

Eingangskennlinie: $I_D = I_{D0} \left(1 - \frac{U_{GS}}{U_{th}} \right)^2 \Rightarrow I_D = \frac{I_{D0}}{U_{th}^2} (U_{GS} - U_{th})^2$

Steilheitskoeffizient: $\beta = 2 \cdot \frac{I_{D0}}{U_{th}^2}$

**Ohne Early-Effekt:
Drainstrom**

Linearer Bereich: $I_D = \beta \left[(U_{GS} - U_{th}) \cdot U_{DS} - \frac{(U_{DS})^2}{2} \right]$

(ohmscher Bereich)

Sättigungsbereich: $I_D = \frac{1}{2} \beta \cdot (U_{GS} - U_{th})^2$

(Arbeitsbereich)

Steilheit: $S = \frac{\partial I_D}{\partial U_{GS}} = \beta (U_{GS} - U_{th})$

**Mit Early-Effekt:
Drainstrom:**

Linearer Bereich: $I_D = \beta \left[(U_{GS} - U_{th}) \cdot U_{DS} - \frac{(U_{DS})^2}{2} \right] \cdot \left(1 + \frac{U_{DS}}{U_A} \right)$

(ohmscher Bereich)

Sättigungsbereich: $I_D = \frac{1}{2} \beta \cdot (U_{GS} - U_{th})^2 \cdot \left(1 + \frac{U_{DS}}{U_A} \right)$

(Arbeitsbereich)

Steilheit: $S = \frac{\partial I_D}{\partial U_{GS}} = \beta (U_{GS} - U_{th}) \cdot \left(1 + \frac{U_{DS}}{U_A} \right)$

Selbstsperrende Isolierschicht- Feldeffekttransistoren:

Steilheitskoeffizient: $\beta = \mu_n C'_{ox} \frac{W}{l} = \mu_n \cdot \frac{\epsilon_0 \cdot \epsilon_{ox}}{t_{ox}} \cdot \frac{W}{l}$

Alle Feldeffekttransistoren:**Source-Schaltung:**

Kleinsignal-Spannungsverstärkung: $A = \frac{u_a}{u_e} \bigg|_{i_a=0} = -S(R_D \parallel r_{DS}) \stackrel{r_{DS} \gg R_D}{\approx} -SR_D$

Kleinsignal-Eingangswiderstand: $r_e = \frac{u_e}{i_e} = \infty$

Kleinsignal-Ausgangswiderstand: $r_a = \frac{u_a}{i_a} = R_D \parallel r_{DS} \stackrel{r_{DS} \gg R_D}{\approx} R_D$

Source-Schaltung mit Stromgegenkopplung:

Kleinsignal-Spannungsverstärkung: $A = \frac{u_a}{u_e} \bigg|_{i_a=0} = -\frac{SR_D}{1 + SR_S} \stackrel{SR_S \gg 1}{\approx} -\frac{R_D}{R_S}$

Kleinsignal-Eingangswiderstand: $r_e = \infty$

Kleinsignal-Ausgangswiderstand: $r_a = \frac{u_a}{i_a} \approx R_D$

Drain-Schaltung:

Kleinsignal-Spannungsverstärkung: $A = \frac{u_a}{u_e} \bigg|_{i_a=0} \approx \frac{SR_S}{1 + (S + S_B)R_S} \stackrel{u_{BS}=0}{=} \frac{SR_S}{1 + SR_S}$

Kleinsignal-Eingangswiderstand: $r_e = \frac{u_e}{i_e} \bigg|_{i_a=0} = \infty$

Kleinsignal-Ausgangswiderstand: $r_a = \frac{u_a}{i_a} \approx \frac{1}{S} \parallel \frac{1}{S_B} \parallel R_S \stackrel{u_{BS}=0}{=} \frac{1}{S} \parallel R_S$

Gate-Schaltung:

Kleinsignal-Spannungsverstärkung: $A = \frac{u_a}{u_e} \bigg|_{i_a=0} \approx (S + S_B)R_D \stackrel{u_{BS}=0}{=} SR_D$

Kleinsignal-Eingangswiderstand: $r_e = \frac{u_e}{i_e} \bigg|_{i_a=0} \approx \frac{1}{S + S_B} \stackrel{u_{BS}=0}{=} \frac{1}{S}$

Kleinsignal-Ausgangswiderstand: $r_a = \frac{u_a}{i_a} \approx R_D$

Differenzverstärker

$$A_G \approx -\frac{R_C}{2R_E}$$

$$A_D \approx +\frac{1}{2}SR_C$$

Operationsverstärker

e-Funktionsgenerator: $u_a = R_N I_{CS} e^{-\frac{u_e}{U_T}}, u_e < 0$

Logarithmierer: $u_a = -U_T \ln 10 \cdot \log \frac{u_e}{I_{CS} R_1}, u_e > 0$

Instrumentationsverstärker: $u_a = \left(1 + \frac{2R_2}{R_1}\right) \cdot (u_{e2} - u_{e1})$

Kondensator

Ladevorgang eines Kondensators:

Halbwertszeit: $t = \ln(2) \cdot \tau$

Gesamtladezeit: $t = 5 \cdot \tau$