

2 Dioden und Transistoren

Student Group

First Name	Surname	Matrikel Nr.

Table of Contents

2. Dioden und Transistoren	2
Einführendes Beispiel	2
Ziele für den Bipolartransistor	2
Ziele für den Feldeffekt-Transistor	3
2.6 Funktionsprinzip eines (Bipolar-)Transistors	3
Schaltzeichen	4
Korrekte Verschaltung der Transistoren	6
Transistor im Bändermodell	6
Kenngröße, Kennlinien, Kennfelder	8
Merke	9
2.7 Funktionsprinzip eines Feldeffekt-Transistors	9
Metall-Oxid-Halbleiter Feldeffekt-Transistor (MOSFET)	10
Ausgangskennlinienfeld	11
Varianten von Feldeffekt-Transistoren	11
Merke	13
Auslegung von Halbleiter-Elementen	13
Merke	14
2.8 Anwendungen	14
Bipolartransistoren	14
Darlington-Transistor	14
Innenleben eines Operationsverstärkers	14
Feldeffekttransistoren	14
NOT Gatter	14
Verpolschutz	15
Pegelwandler	15
Spannungsverdoppler/-invertierer	15
Vollbrücke / H-Brücke	15
Aufgaben	16
Lernfragen	16

2. Dioden und Transistoren

Eine schöne Einführung ist im [KIT Brückenkurs - 3.3.6 Dioden und Transistoren \(*\)](#) zu finden. Aus dieser Einführung sind einige der folgenden Passagen, Videos und Bilder entnommen.

Eine weitere Einführung finden Sie unter [LEIFIphysik](#).

Einführendes Beispiel

Die Elektronik in PCs, Mobiltelefonen, elektrischen Zahnbürsten und wie allen anderen digitalen Begleitern basiert auf Transitorschaltungen (vgl. [im Herzen eines Computers](#)). In [Grundlagen der Digitaltechnik](#) wurde bereits erläutert, dass alle logischen Schaltungen über konjunktive und disjunktive Normalform auf NAND bzw. NOR Gatter zurückgeführt werden kann. Diese wiederum bestehen aus Transistoren. In der unten stehenden Simulation ist die Struktur eines NAND-Gatters in der aktuellen CMOS-Struktur dargestellt. CMOS deutet hierbei auf die Struktur der Schaltung und der Halbleiter Struktur hin: **Complementary metal-oxide-semiconductor** - eine gegensätzlich ergänzende Schaltung aus Halbleitern der Struktur Metall-Oxid-Halbleiter. Die komplementäre Struktur zeigt sich darin, dass

- vom digitalen Ausgang (\$OUT2\$) zur Masse zwei Transistoren einer Art in Reihe geschaltet sind und
- vom digitalen Ausgang (\$OUT2\$) zur 5V-Versorgung zwei Transistoren einer anderen Art parallel geschaltet sind.

Diese zwei unterschiedliche Arten von MOS-Transistoren und weitere verwendete Arten sollen in diesem Kapitel erklärt werden.

Ziele für den Bipolartransistor

Nach dieser Lektion sollten Sie:

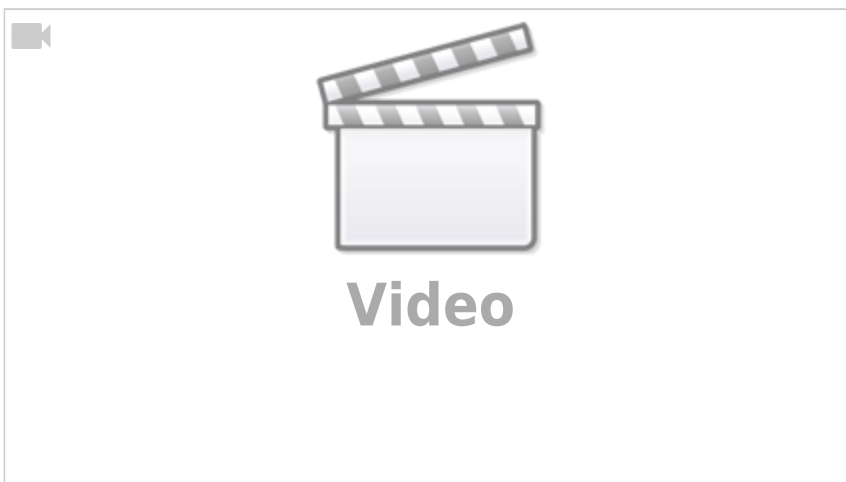
1. wissen, welche Arten von Bipolartransistoren es gibt, wie deren Schichtstruktur und das Schaltsymbol aussieht.
2. wissen, wie die beiden Arten von Bipolartransistoren angesteuert werden.
3. wissen, welche die wichtigsten Kennfelder des Bipolartransistors sind und wie diese aussehen.

Ziele für den Feldeffekt-Transistor

Nach dieser Lektion sollten Sie:

1. wissen, welche Arten von MOSFETs es gibt, wie deren Schichtstruktur und das Schaltsymbol aussieht.
2. wissen, wie das Ausgangskennlinienfeld des MOSFETs aussieht.
3. wissen, was die Bodydiode ist und woher diese herrührt.
4. wissen, was bei der Auslegung eines Halbleiterelements im Ausgangskennlinienfeld zu beachten ist.

2.6 Funktionsprinzip eines (Bipolar-)Transistors



Aus der Diode bzw. dem PN-Übergang heraus lässt sich ein veränderbarer Widerstand entwickeln. Bei diesem gesteuerten Übergangswiderstand ("transfer resistor" oder besser Transistor) kann durch eine Strom der Widerstand verändert und so der durchgelassene Strom eingestellt werden.

[Video-Transkript \(Alternativ zur Erklärung im Video\)](#)

Ein Transistor besteht dabei aus zwei gegeneinander geschalteten Dioden, die eine gemeinsame n- bzw. p-Schicht besitzen, z.B. wird zwischen zwei n-dotierte Schichten eine dünne p-dotierte Schicht gebracht. Dies ist ein npn-Transistor, die häufigere Bauform. Für spezielle Anwendungen werden aber auch pnp-Transistoren eingesetzt. Alle drei Schichten werden elektrisch kontaktiert, der Transistor hat also drei Anschlüsse. Den Kontakt an die mittlere Schicht nennt man Basis (B), die Kontakte zu den beiden äußeren Schichten Kollektor (C) und Emitter (E). Das Schaltzeichen eines Transistors ist in [figure 2](#) dargestellt.

Ein Transistor wird meist als Schalter oder als Stromverstärker betrieben. Um die Funktionsweise zu erläutern, wird in der unten stehenden Abbildung eine typische Transistorschaltung gezeigt. Der Stromkreis, der den Verbraucher, hier die Glühlampe, enthält, nennt man den Arbeitskreis. Hier muss die Spannungsquelle so geschaltet sein, dass die technische Stromrichtung durch den Transistor vom Kollektor zum Emitter verläuft, also in Richtung der am Emitter angezeigten Pfeilrichtung.

Der zweite Kreis, in dem eine positive Steuerspannung an der Basis anliegt, ist der Steuerkreis. Von

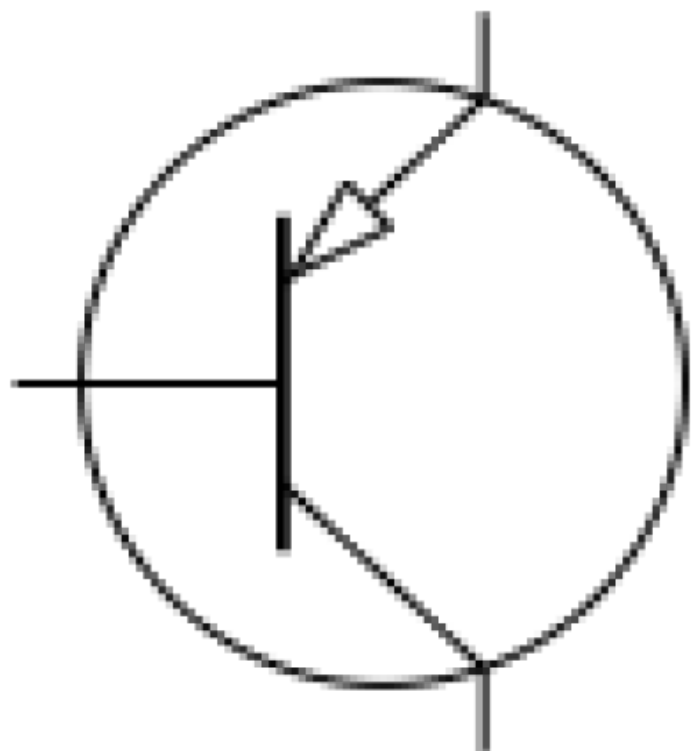
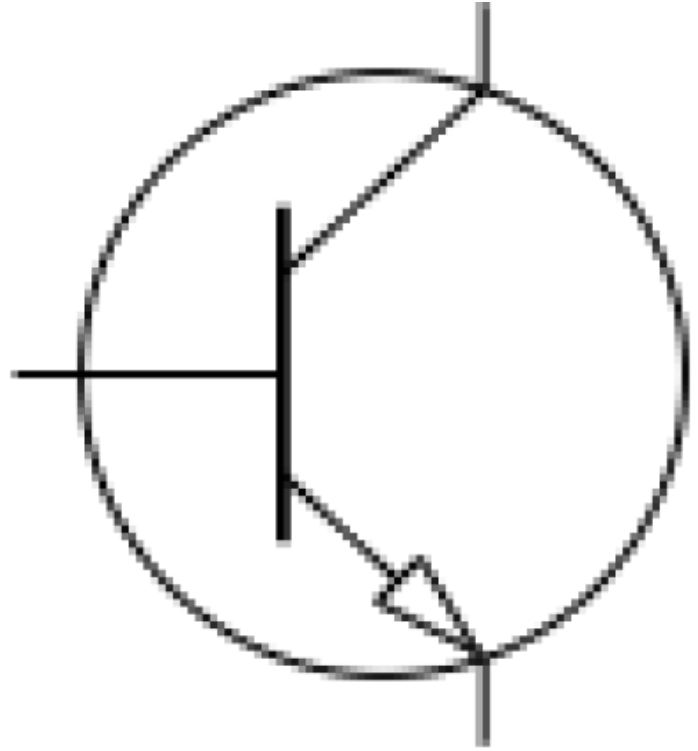
der p-dotierten Basis werden durch die positive Steuerspannung Löcher in den n-leitenden Emitter gepumpt, da an diesem eine negative Spannung anliegt. Die Basis-Emitter-Diode wird also in Durchlassrichtung betrieben. Am Kollektor hingegen liegt eine positive Spannung an, sodass diese Diode eigentlich sperrt.

Überschreitet die Spannung U_{BE} im Steuerkreis eine gewisse Schwelle, kann im Arbeitskreis nun ein Strom I_C fließen. In dieser Hinsicht wirkt der Transistor als Schalter. Mit dem kleinen Strom I_B im Steuerkreis kann daher der große Strom I_C im Arbeitskreis gesteuert werden.

In einem bestimmten Bereich ist der Strom I_B im Arbeitskreis proportional zum Strom I_C im Steuerkreis. Dieses Verhältnis wird als Stromverstärkung $\beta = \frac{I_C}{I_B}$ des Transistors bezeichnet. Verstehen kann man dieses Verhalten, wenn man sich vor Augen führt, dass die p-leitende Basisschicht sehr dünn ist im Vergleich zu den n-leitenden Schichten. Die Elektronen, die über den Steuerkreis zugeführt werden, diffundieren schnell durch sie hindurch und gelangen zu 99% in den mit dem Pluspol verbundenen Kollektor und werden durch den Arbeitskreis in den Emitter zurück gepumpt. Nur wenige gelangen durch den Emitter direkt wieder in den Steuerkreis. Daher ist der Strom im Steuerkreis sehr viel geringer als der Strom im Arbeitskreis.

Schaltzeichen

Fig. 1: Schaltzeichen und vereinfachter Aufbau von npn- und pnp-Bipolartransistors



Wie gerade beschrieben, ist der Bipolartransistor durch eine dreilagige, abwechselnd dotierte Schichtstruktur aufgebaut, welche zwei entgegengesetzten und hintereinander geschalteten Dioden entspricht. Abhängig von der Schichtfolge (bzw. "Richtung der Dioden") ergeben sich pnp- bzw. npn-Transistoren, die durch unterschiedliche Schaltzeichen mit drei Anschlüssen dargestellt werden (siehe [figure 1](#)). Bei beiden Transistorvarianten werden vom Emitteranschluss (E) Ladungsträger Richtung Kollektoranschluss (C) ausgesandt, falls durch den Basisanschluss (B) ein geeigneter Strom fließt. Vereinfacht gesehen, könnten die negativen Ladungsträger der n-dotierten Seiten einen Strom durch eine NPN-Struktur darstellen, wenn auch in der p-dotierten Schicht negativen Ladungsträger vorhanden wären. Der damit fließende Strom I_C in der technischen Stromrichtung wird im Schaltzeichen über die Pfeilrichtung am Emitter verbildlicht. Beim NPN-Transistor fließt also der Strom I_C vom Kollektor zum Emitter. Da beim PNP-Transistor positive Ladungsträger die Leitfähigkeit ermöglichen, zeigt hier die technische Stromrichtung vom Emitter zum Kollektor und der Pfeil am Emitter zum Kollektor hin. Die Pfeilrichtung gleicht der Richtung der Diode bzw. des PN-Übergangs. Weitere Eselsbrücken zur Pfeilrichtung sind:

- „Tut der Pfeil der Basis weh, handelt's sich um **pnp**“
- „Will der Pfeil sich von der Basis trenn', handelt sich's um **nnp**“

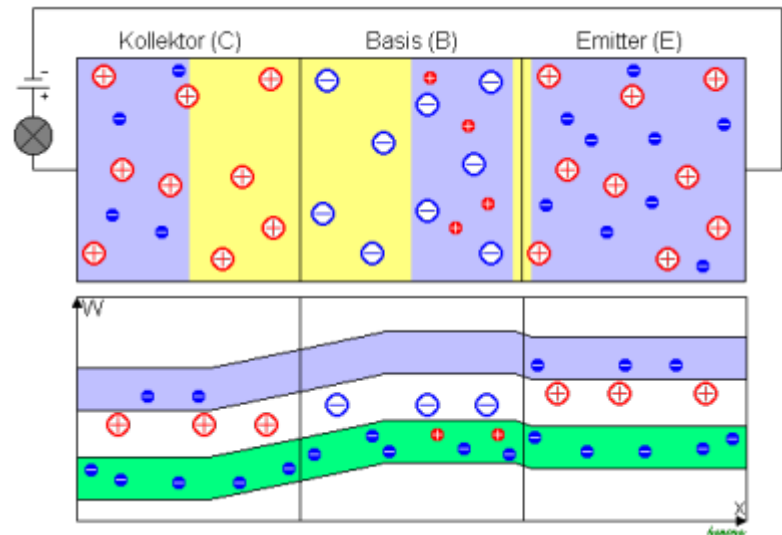
Korrekte Verschaltung der Transistoren

In der Simulation rechts ist die korrekte Verschaltung der Transistoren zu sehen. Allgemein muss bei der korrekten Verschaltung der Pfeil des Symbols der technischen Stromrichtung zeigen. Der Basisstrom I_B wird in den Schaltungen fast immer durch eine Spannungsquelle zwischen Basis und Emitter mit einer Spannung U_{BE} erzeugt. Dabei wird beim NPN-Transistor eine positive Spannung gegenüber des Emitters benötigt und beim PNP-Transistor eine negative Spannung gegenüber des Emitters. In der praktischen Anwendung überwiegen die NPN-Transistoren, unter anderem da die dort genutzten negativen Ladungsträger eine höhere Leitfähigkeit erzeugen. Für die folgenden Erklärungen werden nur NPN-Transistoren betrachtet.

Eine zentrale Frage die sich bei einem näheren Blick auf die Simulation rechts ergibt, ist: Warum muss beim NPN-Transistor ein technischer Strom in die Basis fließen, also positive Ladungsträger in die P-Schicht, zugeführt werden? Wäre es nicht einleuchtender, wenn die nicht vorhandenen und zum Transport benötigten negativen Ladungsträger zugeführt werden müssten?

Transistor im Bändermodell

Fig. 2: Transistor im Bändermodell



Um das zu verstehen, werden die Erkenntnisse des PN-Übergangs benötigt. Im Bild [figure 2](#) ist der Aufbau des NPN-Transistors im Bändermodell zu sehen. Im n-dotierten Kollektor und Emitter sind die frei beweglichen negativen Ladungsträger (blau) und ortsfesten positiven Ladungsträger (rot) eingezeichnet, in der Basis entsprechend die frei beweglichen positiven Ladungsträger (rot) und ortsfesten negativen Ladungsträger (blau). Beide PN-Übergänge haben eine Sperrschicht ausgebildet. Am Transistor liegt eine positive Spannung U_{CE} an, welche in der dargestellten Situation keinen Stromfluss erzeugen kann. Durch die positive Spannung U_{CE} und dem fehlenden Potential an der Basis sinkt die Spannung U_{BE} , was zu einer Verkleinerung der Sperrschicht führt. Im Gegensatz dazu erhöht sich die Spannung $U_{CB} = U_{CE} - U_{BE}$. Damit wird die Sperrschicht zwischen Basis und Kollektor größer. Bei Variation der äußeren Spannung U_{CE} wird immer mindestens ein PN-Übergang vorhanden sein, der in Sperrrichtung verschaltet ist, das heißt der Transistor sperrt.

Um die Sperrschicht zwischen Kollektor und Basis aufzuheben, muss diese in Durchlassrichtung verschaltet werden. Bis zum Durchschalten des Transistors sind dies mehrere Schritte:

1. Zunächst ermöglicht eine positive Spannung U_{BE} das Durchschalten des PN-Übergangs zwischen Basis und Emitter. Dazu werden durch den Strom I_B Löcher in der Basis bereitgestellt. In der Simulation rechts ist zu sehen, dass die Verschaltung des Transistors in der Art ist, dass in der (physikalisch nicht ganz korrekten) Diodenschaltung die Diode zwischen Basis und Emitter leitend wird.
2. Durch den so fließenden Strom vom Emitter sind Elektronen in der Basis vorhanden. Da im realen System die Basis nur wenige Mikrometer groß ist, findet dort so gut wie keine Rekombination der Elektronen mit den Löchern statt.
3. Die vorhandenen Löcher und Elektronen schwächen die Sperrschicht zwischen Basis und Kollektor, sodass diese aufgehoben wird.
4. Die Elektronen aus dem Emitter können nun die Basis durchschreiten. Die Basis ist im Vergleich zur mittleren freien Weglänge der Ladungsträger ("Weg bis zur Rekombination mit einem Loch") sehr klein.
5. Beim NPN-Bipolartransistor tragen also sowohl Löcher (zum Aufheben der Sperrschichten), als auch Elektronen (als "Hauptverantwortliche" für den Ladungstransport, die sog. Majoritätsträgerladungen) zur Leitfähigkeit bei. Daher rührt der Name Bipolartransistor

Kenngröße, Kennlinien, Kennfelder

Im vorherigen Kapitel wurde bereits schon auf Kenngröße einer Blackbox eingegangen, dort speziell für einen Verstärker. Die Methodik kann hier auch angewandt werden. Im Video oben wurde bereits schon die erste Kenngröße beschrieben: Die **Stromverstärkung** $\beta = \frac{d I_C}{d I_B}$, bzw. in Form einer Grafik, die **Stromsteuerkennlinie** $I_C(I_B)$.¹⁾

Eine weitere Kennlinie ist das **Eingangskennlinienfeld** $U_{BE}(I_B)$ bzw. als differentielle Kenngröße (=Steigung in der Kennlinie) der **differentielle Eingangswiderstand** $r_{BE} = \frac{d U_{BE}}{d I_B}$. Wie bereits beschrieben, gleicht der Aufbau zwischen Basis und Emitter einer Diode. Das Eingangskennlinienfeld gleicht dementsprechend dem einer Diode. Da der Stromfluss I_B sehr klein ist (wenige Mikroampere oder kleiner), ist der Eingangswiderstand r_{BE} groß.

Die folgende Simulation zeigt die Stromsteuerkennlinie $I_C(I_B)$ und Eingangskennlinie(nfeld) $U_{BE}(I_B)$ durch Variation von U_{BE} (bzw. I_B).

Für die Beschreibung des Transistors ist das **Ausgangskennlinienfeld** $U_{CE}(I_C)$ und der darin als Steigung vorhandene **differentielle Kollektor-Emitter-Widerstand** $r_{CE} = \frac{d U_{CE}}{d I_C}$ besonders wichtig. Dieses ist in der folgenden Simulation für unterschiedliche Eingangsspannungen U_{BE} (und damit unterschiedlichen Steuerströmen I_B) zu sehen. Das Ausgangskennlinienfeld lässt sich in verschiedene Bereiche unterteilen:

1. **Sperrbereich:** bei geringen Eingangsspannungen $U_{BE} < 600\text{mV}$ wird die Sperrschicht nicht abgebaut. Entsprechend wird der gesamte Transistor nicht leitend. Im Ausgangskennlinienfeld ist dies dadurch zu sehen, dass bei positiver Ausgangsspannung U_{CE} der Ausgangsstrom I_C sehr klein wird. In diesem Fall entspricht der Transistor auf der Ausgangsseite einem hochohmigen Widerstand, bzw. einem offenen Schalter.
2. **Verstärkungsbereich** (oder aktiver Bereich): bei größeren Eingangsspannungen $U_{BE} > 600\text{mV}$ wird die Sperrschicht abgebaut. Im Verstärkungsbereich verhält sich die Ausgangskennlinie wie eine Gerade. Der Ausgangsstrom I_C ist damit nur noch abhängig von I_B , so wie dies über die Stromverstärkung $\beta = I_C/I_B$ definiert ist.
3. **Sättigungsbereich:** Der Sättigungsbereich ist bei größeren Eingangsspannungen $U_{BE} > 600\text{mV}$ und nur geringer Ausgangsspannung U_{CE} zu finden. Bei konstanter Eingangsspannung U_{BE} verhält sich die Ausgangsspannung zum Ausgangsstrom wie ein hoher, nicht-linearer Widerstand. In diesem Fall entspricht der Transistor auf der Ausgangsseite einem niederohmigen Widerstand, bzw. einem leitenden Schalter.

Im Datenblatt ist gelegentlich eine andere Nomenklatur zu finden, die sich aus der sogenannten **H-Charakteristik der Vierpoltheorie**²⁾ ergibt:

- Stromverstärkung $h_{fe} = \beta(I_C, U_{CE}) = \frac{I_C}{I_B}$
- Eingangswiderstand $h_{ie} = r_{BE}(I_C, U_{CE}) = \frac{U_{BE}}{I_B}$
- Ausgangswiderstand $h_{oe} = r_{CE}(I_B, U_{BE}) = \frac{U_{CE}}{I_C}$

Der Bipolartransistor wird dort genutzt, wo eine geringe Schwellspannung oder ein Stromverstärker benötigt wird. Die ist beispielsweise bei verschiedenen Verstärkerschaltungen vorteilhaft. Auch in einigen einfachen Netzteilen sind Bipolartransistoren zu finden. Die häufigste

Bipolartransistorschaltung ist die sogenannte Kollektorschaltung. Diese ist dadurch gekennzeichnet, dass am Kollektor eine konstante Spannung - die Versorgungsspannung - anliegt. Mehrere Kollektorschaltungen können durch eine gemeinsamen Spannungsversorgung betrieben werden. Damit liegt an allen Kollektoranschlüssen die gleiche Spannung an. Aufgrund der breiten Verwendung der Bipolartransistoren hatten, wird auch heute noch die gemeinsame Spannungsversorgung von elektronischen Schaltungen V_{CC} genannt, wobei CC für **Common Collector** steht. Dies ist häufig selbst dann noch zu sehen, wenn keine Bipolartransistoren mehr verwendet werden.

Ein großer Nachteil des Bipolartransistors ist, dass für das Schalten ein Steuerstrom benötigt wird. Besonders in digitalen Schaltungen, aber auch in der Leistungselektronik, ergibt sich dadurch eine nicht zu vernachlässigende Eingangsleistung $P = U_{BE} \cdot I_B$. Diese führt zu Verlusten und Abwärme, die bei der Leistungsversorgung und thermischen Auslegung berücksichtigt werden müssen. Aus diesem Grund werden in aktuellen Mikrocontrollern keine Bipolartransistoren mehr genutzt. In diesen Feldern wurde der Bipolartransistor durch den Feldeffekt-Transistor verdrängt.

Merke

Es gibt 2 verschiedene Arten von Bipolartransistoren. Diese unterscheiden sich in der Art der Schichtenaufbaus, bzw. der Majoritätsträgerladungen:

- **npn-Bipolartransistoren:** Die hauptsächliche Leitung geschieht über Elektronen. Diese können ohne Strom I_B über die Basis den p-dotierten Bereich nicht durchlaufen. Durch $I_B < 0$ werden Löchern in die Basis eingebracht, welche Sperrschichten abbauen.
- **pnp-Bipolartransistoren:** Die hauptsächliche Leitung geschieht über Löcher. Diese können ohne Strom I_B über die Basis den n-dotierten Bereich nicht durchlaufen. Durch $I_B > 0$ werden Elektronen in die Basis eingebracht, welche Sperrschichten abbauen.

Beim Bipolartransistor sind beide Ladungsträgertypen am Transport beteiligt.

2.7 Funktionsprinzip eines Feldeffekt-Transistors

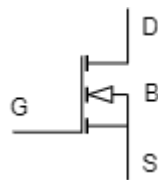


Fig. 6: FET Schaltsymbole

Auch ein Feldeffekt-Transistor (FET) besteht aus zwei gegeneinander geschalteten Dioden, die eine gemeinsame n- bzw. p-Schicht besitzen. Die Leitfähigkeit des Feldeffekt-Transistors wird jedoch nicht durch das Anlegen einer Steuerstroms, sondern allein durch eine Steuerspannung erzeugt. Auch beim Bipolartransistor wurde der Steuerstrom durch eine Steuerspannung generiert. Der Steuerstrom muss jedoch zum Ansteuern des Bipolartransistors dauerhaft fließen, da die über die Basis eingetragenen Ladungsträger intern rekombinieren.

In [figure 6](#) ist ein spezieller Feldeffekt-Transistor gezeichnet, dem sogenannten "Metall-Oxid-Halbleiter

Feldeffekt-Transistor". Dieser soll im Folgenden näher erklärt werden.

Um die Transistortypen zu unterscheiden, und die dahinterliegende Physik zu betonen, werden die Anschlüsse beim Feldeffekt-Transistor anders bezeichnet:

- **(S) Source:** Anschluss, von dem die Ladungsträger aus den Transistor durchlaufen (entspricht etwa dem Emitter)
- **(G) Gate:** Anschluss, an dem mittels einer Spannung die Leitfähigkeit geändert werden kann (entspricht etwa der Basis, wobei dort Steuerströme eingespeist werden)
- **(D) Drain:** Anschluss, an dem die Ladungsträger ankommen und aus dem Transistor austreten (entspricht etwa dem Kollektor)

Daneben gibt es im Aufbau noch das **"Bulk" (B)**, was das Grundsubstrat des Transistors bezeichnet. Dies ist in der Regel nicht separat herausgeführt, sondern mit dem Sourceanschluss kurzgeschlossen. Bei einigen FETs ist das Bulk durch die mittlere Verbindung dargestellt.

In der Simulation rechts ist zu sehen, dass sich der Feldeffekt-Transistor so ähnlich wie ein Schalter verhält, welcher über eine Spannung gesteuert wird. Auf das Gate scheint kein Strom zu fließen, aber wenn sich die Spannung am Gate ändert, so ändert sich das Verhalten von "leitfähig" in "offen".

Metall-Oxid-Halbleiter Feldeffekt-Transistor (MOSFET)

Fig. 7: MOSFET Schichtung 

Der Aufbau des Metall-Oxid-Halbleiter Feldeffekt-Transistors (englisch **Metal Oxide Semiconductor Field Effect Transistor: MOSFET**) ähnelt auf dem ersten Blick dem Bipolartransistor. In [figure 7](#) ist in den einzelnen Bildern (1)...(3) die Schichtung eines n-Kanal (englisch n-Channel) MOSFETs und in (4) nochmals das Schaltsymbol dargestellt. Im Gegensatz zum npn-Bipolartransistor ist hier aber mittlere p-dotierte Schicht (Bulk) nicht direkt an der Steuerelektrode angeschlossen. Vielmehr bildet die Metallschicht des Gates ([figure 7](#), Bild (5), grau), die Isolationsschicht des Oxids (im Bild violett) und die leitfähige, p-dotierte Schicht des Bulks (im Bild rot) einen Kondensator. Dabei ist zu beachten, dass das Bulk auf dem Potential des Source-Anschlusses liegt (gestichelte Linie im Bild).

Ohne Spannungsdifferenz U_{GS} zwischen Gate und Source bildet sich an den p-n Übergängen eine (kleine) Sperrschicht aus. Wird die Spannungsdifferenz U_{GS} vergrößert, so wird der Kondensator zwischen Gate und Bulk aufgeladen. Damit reichern sich gegenüber der Gatelektrode Elektronen an ([figure 7](#), Bild (2), dunkelblauer "Keil"). Übersteigt die Spannungsdifferenz U_{GS} eine bestimmte Schwellspannung, so bilden die angereicherten Elektronen einen Kanal zwischen Source und Gate. Damit kann ein Strom $I_D \gg 0$ durch den MOSFET ([figure 7](#) Bild (3)) fließen.

Das Schaltsymbol ([figure 7](#), Bild (4)) lässt sich auch folgendermaßen beschreiben: Es bilden sich im ausgeschalteten Zustand jeweils Kondensatoren zwischen Gate und Source, zwischen Gate und Basis und zwischen Gate und Drain wegen der Oxidschicht (im Bild (1) violett) aus.³⁾ Um den MOSFET anzusteuern, muss die Spannung am Gate U_{GS} so geartet sein, dass sich ein PN-Übergang im Bulk bildet, angedeutet durch das weiß ausgefüllte Dreieck im Bild (4). Da die Spitze des Dreiecks (bzw. des damit skizzierten Diodensymbols) in Richtung des Gates zeigt, ist klar, dass es sich um einen n-Kanal MOSFET handelt.

In der Simulation rechts sind die gleichen Spannungsverhältnisse wie in [figure 7 \(1\)...](#)(3) dargestellt. Durch den Wechselschalter links ist es möglich die Spannung U_{DS} über den Transistor zu invertieren. Wird diese negativ, so stellt sich eine etwas andere Situation ein: Der MOSFET scheint leitfähig zu werden, unabhängig davon, welche Spannung U_{GS} annimmt. Dies ist darauf zurückzuführen, dass sich im Schichtenaufbau eine weitere Diode versteckt hat: zwischen Bulk (p) und Drain (n) hat sich eine Sperrschicht ausgebildet, welche bei $U_{DS} < 0$ und mit der Anbindung von Bulk und Source in Durchlassrichtung betrieben wird. Diese sogenannte Body-Diode ist in der Simulation bei (3b) explizit eingebaut.

Ausgangskennlinienfeld

Auch beim MOSFET soll das **Ausgangskennlinienfeld** $U_{DS}(I_D)$ betrachtet werden. Auch dieses ähnelt dem Bipolartransistor, jedoch sind nun die verschiedenen Kennlinien durch unterschiedliche Steuerspannungen U_{GS} und nicht durch einen Steuerstrom einstellbar.

Leider unterscheidet sich die Benennung der verschiedenen Betriebsbereiche eines MOSFETs von denen des Bipolartransistors:

1. Sperrbereich: bei geringen Eingangsspannungen U_{GS} kann kein Kanal gebildet werden. Entsprechend wird der gesamte Transistor nicht leitend. Im Ausgangskennlinienfeld ist dies dadurch zu sehen, dass bei positiver Ausgangsspannung U_{DS} der Ausgangsstrom I_D sehr klein wird. In diesem Fall entspricht der Transistor auf der Ausgangsseite einem hochohmigen Widerstand, bzw. einem offenen Schalter.
2. Sättigungsbereich: bei größeren Eingangsspannungen $U_{GS} > U_{th}$ oberhalb eines Grenzwerts (engl. Threshold) wird ein leitfähiger Kanal gebildet. Im Sättigungsbereich verhält sich die Ausgangskennlinie wie eine Gerade. Der Ausgangsstrom I_D ist damit nur noch abhängig von U_{GS} .
3. linearer Bereich (aktiver Bereich) : Der lineare Bereich ist bei größeren Eingangsspannungen $U_{GS} > U_{th}$ und nur geringer Ausgangsspannung U_{DS} zu finden. Bei konstanter Eingangsspannung U_{GS} verhält sich die Ausgangsspannung zum Ausgangsstrom wie ein hoher, nicht-linearer Widerstand. In diesem Fall entspricht der Transistor auf der Ausgangsseite einem niederohmigen Widerstand, bzw. einem leitenden Schalter.

Zu beachten ist, dass der Sättigungsbereich bei MOSFET und Bipolartransistor unterschiedliche Arbeitsbereiche kennzeichnet.

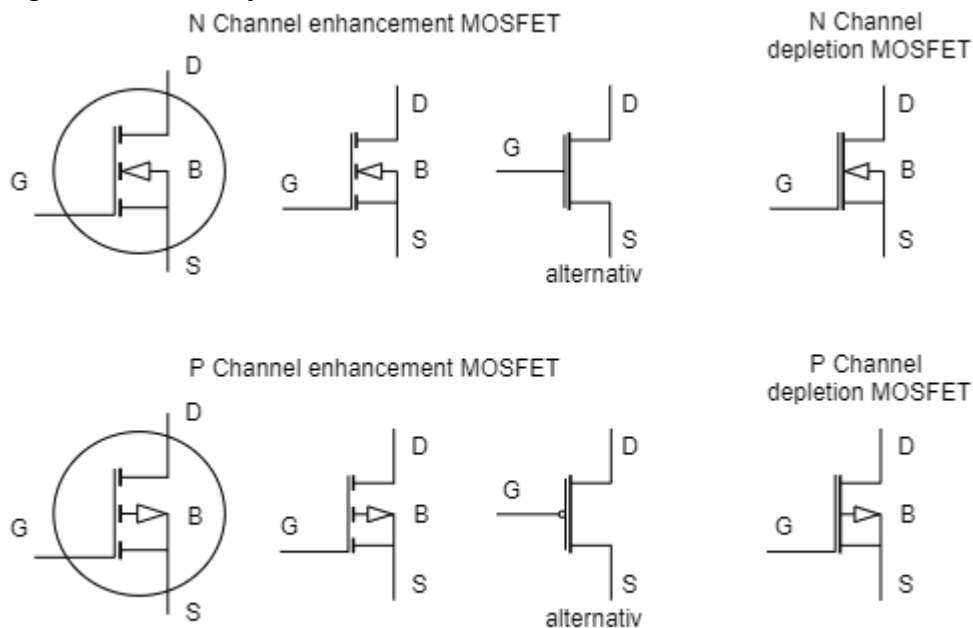
Varianten von Feldeffekt-Transistoren

Der bisher betrachtete (und auch am häufigsten genutzte) Feldeffekt-Transistor ist der sogenannte **“n-Kanal Anreicherungstyp MOSFET”**. Dabei rührt der Teil “n-Kanal” vom Typ des Strom-bildenden Ladungsträgers und wurde bereits weiter oben gegeben. Der Teil “Anreicherungstyp” (engl. enhancement) stellt dar, dass die Ladungsträger zunächst nicht vorhanden sind und zur Leitfähigkeit

erst mittels der der Spannung U_{GS} im Bulk angehuft werden mssen.

Bei einigen Schaltungen (insbesondere Digitalschaltungen) werden auch "**p-Kanal Anreicherungstyp MOSFET**" verwendet, bei dem Locher die Strom-bildenden Ladungstrger sind. In der Simulation rechts ist diese Art des MOSFET gezeigt. Am deutlichsten ist, dass bei der Verschaltung des p-Kanal Anreicherungstyp MOSFET in der Regel Drain und Source vertauscht wird. Damit werden die Zahlenwerte von U_{DS} und I_D im Ausgangskennlinienfeld negativ. Um Locher im p-Kanal anzureichern muss eine negative Spannung am Gate $U_{GS} < 0$ anliegen.

Fig. 8: FET Schaltsymbole



CC BY-SA 4.0 Hochschule Heilbronn - Mechanik und Robotik

In der [figure 8](#) sind die Schaltsymbole verschiedener Varianten von MOSFETs dargestellt. In den MOSFETs der oberen Zeile wird ein n-Kanal zum Ladungstransport ausgebildet, in der unteren ein p-Kanal.

In [figure 8](#) links oben sind drei Varianten eines **n-Kanal Anreicherungstyp MOSFET** gezeigt. Beim ersten Schaltsymbol stellt der Kreis dar, dass es sich um ein diskretes Bauteil handelt, also ein einzelnen MOSFET, der nicht mit anderen gemeinsam in einem Chip integriert ist. Das zweite Schaltsymbol wurde bereits in den vorherigen Kapitel genutzt. Das dritte Schaltsymbol des gleichen n-Kanal Anreicherungstyp MOSFET ist die reduzierte Variante (also ohne Bulk). Diese Darstellung wird zur Vereinfachung in Digitalschaltungen genutzt.

In [figure 8](#) links unten sind drei Varianten eines **p-Kanal Anreicherungstyp MOSFET** gezeigt. Auch hier zeigt der Kreis beim ersten Schaltsymbol an, dass es ein diskretes Bauteil ist, jedoch ist nun die die Pfeilrichtung am Bulk gedreht. Das zweite Schaltsymbol wird - wie beim n-Kanal MOSFET - so auch in integrierten Schaltungen verwendet. Das dritte Schaltsymbol ist wieder die reduzierte Variante (also ohne Bulk). Fur die Digitalschaltung ist nur wichtig, ob der Schalter bei High-Signal ($V = 5V$) schliet oder ffnet. Da der p-Kanal Anreicherungstyp MOSFET ffnet, wird dieser mit einem Negierungszeichen (kleiner Kreis) am Gate gezeichnet.

In [figure 8](#) rechts sind sogenannte **n-Kanal und p-Kanal Verarmungstyp MOSFET** dargestellt. Die

bisherig betrachteten MOSFETs waren im ausgeschalteten Zustand (also $U_{DS}=0$) nicht leitfähig. In manchen Anwendungen wäre es aber gut, wenn der MOSFET im ausgeschalteten Zustand einem leitfähigen Schalter gleicht. Mit Blick auf die Schichtstruktur (figure 7, Bild (1)...(3)) ist dies über eine gezielte Umdotierung des Bereichs gegenüber des Gates möglich. Durch die Dotierung kann ein leitfähiger Kanal verstellt werden. Die Ladungsträger dieses Kanals können durch ein geeignetes Feld - und damit geeignete Gatespannung U_{GS} - verdrängt bzw. verarmt werden. Damit wird der MOSFET bei einer Gegenspannung U_{GS} nicht-leitend. Im Schaltsymbol ist der "Kurzschluss" zwischen Source und Drain auch bildlich eingezeichnet.

Merke

Es gibt 4 verschiedene Arten von MOSFETs. Diese unterscheiden sich einerseits in der Art der Strom-bildenden Ladungsträger:

- **n-Kanal:** Die Strom-bildenden Ladungsträger sind Elektronen.
- **p-Kanal:** Die Strom-bildenden Ladungsträger sind Löcher.

Das zweite Unterscheidungsmerkmal ist die Leitfähigkeit im ausgeschalteten Zustand ($U_{GS}=0$):

- **Anreicherungstyp:** Bei einer Gatespannung von $U_{GS}=0$, ist kein leitfähiger Kanal vorhanden. Erst durch das Aufladen des Gate-Bulk-Kondensators wird der Kanal gebildet bzw. die Ladungsträger angereichert.
- **Verarmungstyp:** Bei einer Gatespannung von $U_{GS}=0$, ist ein leitfähiger Kanal vorhanden. Durch das Aufladen des Gate-Bulk-Kondensators wird der Kanal verkleinert bzw. die Ladungsträger verdrängt ("verarmt").

Beim Feldeffekttransistor werden durch das elektrische Feld des Gate-Bulk-Kondensators nur genau diejenigen Ladungsträger angereichert bzw. verarmt, welche zum Ladungstransport beitragen.

Fig. 9: Arbeitsbereich von Halbleiter-Elementen 

Auslegung von Halbleiter-Elementen

Bei allen Transistoren und Dioden sind für die Schaltungsauslegung verschiedene die Grenzwerte zu beachten. Diese können im Ausgangskennlinienfeld direkt eingetragen werden (figure 9, oben). Durch die Erwärmung des Bauteils und der daraus steigenden Eigenleitung ergeben sich zwei Grenzwerte:

- Im leitenden Zustand bildet die Verlustleistung $P_{loss}=R(T) \cdot I^2$ einen direkten Bezug zum Strom durch das Halbleiterelement I_C, I_D, I_D (Bipolartransistor, MOSFET, Diode). Daraus ergibt sich Strom ein I_{max} , welcher nicht überschritten werden soll.
- Im Zustand, bei dem sowohl ein merklicher Strom als auch eine merkliche Spannung anliegt, ergibt sich eine maximal erlaubte Leistung $P_{tot}=const.=U \cdot I$. Dies ist in der Ausgangskennlinie eine Hyperbel. Überschreitet der Ausgangsstrom diese Hyperbel, so erwärmt sich das Halbleiterelement so stark, dass durch die ansteigende Eigenleitung, die Leitfähigkeit

absinkt, was wiederum zu einem steigenden Strom führt. Dieser Effekt führt zur thermischen Zerstörung der Komponente.

Daneben darf eine maximale Spannung $U_{\max}$ nicht überschritten werden. Diese ist in meist auf die (interne) Durchschlagsfestigkeit des Bauteils zurückzuführen.

Diese Grenzen sind insbesondere dann wichtig, wenn z.B. ein MOSFET als Schalter genutzt werden soll (Beispiel: [figure 9](#), unten). In diesem Fall gibt es zwei Zustände:

- Schalter ist leitfähig: Es liegt eine geringe Spannung U_{DS} an, bei dem ein großer Strom $I_D < I_{\max}$ fließt.
- Schalter ist nicht leitfähig: Es liegt eine hohe Spannung $U_{DS} < U_{\max}$ an, bei dem kein Strom fließt.

Beim Umschalten von "leitfähig" nach "nicht leitfähig" kann dabei - selbst, wenn die einzelnen Strom- und Spannungsgrenzen berücksichtigt werden - der Schalter zerstört werden. In [figure 9](#) ist dieser Fall im unteren Diagramm zu sehen. Der Stromfluss I_D wird zunächst aufrecht erhalten (bzw. sind nur gering), obwohl Spannung U_{DS} ansteigt (blaue Linie). In diesem Fall kann es dazu kommen, dass P_{tot} überschritten wird und der MOSFET wegen thermischer Überlastung zerstört wird.

Merke

Bei jedem Halbleiterelement sind am Ausgang drei Maximalwerte zu beachten:

- einer maximalen Spannungsgrenze $U_{\max}$,
- einer maximalen Stromgrenze $I_{\max}$,
- einer maximalen Leistungsgrenze $P_{\text{tot}} = U \cdot I$

2.8 Anwendungen

Bipolartransistoren

Darlington-Transistor

Innenleben eines Operationsverstärkers

Feldeffekttransistoren

- TFT-Bildschirme

NOT Gatter

Verpolschutz

Pegelwandler

Der Pegelwandler (auch Logic Level Converter, Level Shifter) ermöglicht die bidirektionale Verbindung von digitalen Anschlüssen unterschiedlicher Spannungsniveaus, z.B. 5 V auf 3,3 V.

Für den Levelconverter kann jeder n-Kanal enhancement MOSFET genutzt werden, dessen Schwellspannung unter 1,8...2 V liegt. Diese Grenze ist auf den Mindestlogikpegel von 2,0 V für logisch 1 zurückzuführen. Der Einfachheit halber werden "logic level enhancement mode MOSFET" eingesetzt, welche gerade für die Logikspannung von 3.3V optimiert sind.

Die Funktionsweise ist auf [Wikipedia](#) gut erklärt und kann sich mit der Simulation hergeleitet werden.

Spannungsverdoppler/-invertierer

Im Oszilloskop wird die Spannung U_{C1} am Eingangskondensator C1 und U_{C2} am Speicherkondensator C1 angezeigt.

Lernfragen:

- In welchem Zustand ist die Spannung U_{C1} gleich 1 V?
In welchem Zustand ist die Differenz der Spannungen $U_{C2} - U_{C1}$ an den beiden Kondensatoren gleich 1V?
- Was passiert, wenn die Spannungsquellen für 0V und 1V vertauscht werden?
- Wie lässt sich diese Schaltung statt mit Wechselschaltern mit Dioden realisieren?

Spannungsinvertierer im Mikrocontroller

weiterführende Informationen unter [Inside the 8087's substrate bias circuit](#)

Vollbrücke / H-Brücke

Aufgaben

Lernfragen

- Beschreiben sie die Funktion eines Transistors
 - Skizzieren Sie den Schichtenaufbau eines Bipolartransistors. Erklären Sie das Durchschalten eines PNP-Bipolartransistors mit Hilfe der gezeichneten Skizze.
 - Zeichnen sie das vereinfachte Dioden-Ersatzschaltbild eines NPN Transistors beschreiben Sie die Funktionsweise.
- Erklären Sie den Unterschied zwischen einem PNP und NPN Transistor.
 - Zeichnen Sie jeweils eine Schaltung, bei dem der jeweilige Schalter an $U_+ = 5V$ und Masse so angeschlossen ist, dass ein Durchschalten mit einer Spannung zwischen U_+ und Masse an der Basis möglich ist.
 - Benennen Sie die jeweiligen Anschlüsse der Transistoren in der Zeichnung.
 - Welche Spannung muss jeweils an der Basis angelegt werden, damit der Transistor durchschaltet?
 - Wie ist jeweils das Vorzeichen des Steuerstroms zu wählen?
 - In welchem Größenbereich liegt eine typische Stromverstärkung?
- stromgesteuerte und spannungsgesteuerte Transistoren
 - Erläutern Sie den Unterschied zwischen einem stromgesteuerten und einem spannungsgesteuerten Transistor.
 - Welche Transistorart ist stromgesteuert, welche spannungsgesteuert?
 - Zeichnen Sie jeweils ein Schaltzeichen für einen stromgesteuerten und einen spannungsgesteuerten Transistor.
 - Wie ist die Reihenfolge der Dotierung der gezeichneten Transistoren?
- Welche zwei grundlegenden Typen von Transistoren gibt es?
- MOSFET
 - Welche Vorteile hat ein MOSFET gegenüber einem Bipolartransistor?
 - Wie ist ein MOSFET aufgebaut? (Schichtstruktur, Anschlüsse)
- H-Brücke
 - Zeichnen Sie eine H-Brücke mit Schaltern (idealer Schalter), einer ohmsch/induktiven Last und einer externen Spannungsquelle mit V_+ und GND.
 - Wie können die verschiedenen Schalter angesteuert werden, um eine beliebige Spannung zwischen V_+ und V_- an der Last anliegen zu haben? Was ist der Fachausdruck für die Ansteuerungsart?
- Zeichnen Sie das notwendige PWM-Signal, um bei einer Vollbrücke eine sinusförmige Ausgabe zu generieren.
- Verwendungsmöglichkeiten für Transistoren
 - Welche Verwendungsmöglichkeiten gibt es für Transistoren?
 - Zeichnen Sie einen Spannungsverdoppler.
 - Was ist ein Pegelwandler?
 - Warum werden heutzutage bevorzugt Feldeffekt-Transistoren und nicht Bipolartransistoren verwendet?

Referenzen zu den genutzten Medien

Element	Lizenz	Link
Video: Stromkreiselemente - Dioden und Transistoren - Teil 4	CC-BY (Youtube)	https://www.youtube.com/watch?v=KjyHta5p9WE

<--

¹⁾ In der Praxis wird noch zwischen der Kleinsignal-Stromverstärkung $\beta = h_{fe}$ und Großsignal-Stromverstärkung $B = h_{FE}$ unterschieden. Beim Kleinsignalverhalten wird eine relativ kleine Änderung um einen festen Arbeitspunkt (z.B. um bestimmte Werte I_C und U_{CE}) betrachtet. Beim Großsignalverhalten wird eine Änderung zwischen 0 und einem gegebenen Wert betrachtet. Bei nichtlinearen Kenngrößen können sich beide Größen unterscheiden. In diesem Kurs wird nur das Kleinsignalverhalten beschrieben. Das Großsignalverhalten und die Unterscheidung beider Betrachtungen wird in diesem Kurs nicht betrachtet.

^{2), 3)} Bei den Feldeffekt-Transistoren bildet sich zusätzlich ein Kondensator zwischen Source und Drain aus, welcher insbesondere beim schnellen Schalten von Induktivitäten zu Überspannungen am MOSFET führen kann.

From:

<https://first.mexle.te.hs-heilbronn.de/> - MEXLE Wiki

Permanent link:

https://first.mexle.te.hs-heilbronn.de/elektronische_schaltungstechnik/2_transistoren?rev=1587392068

Last update: 2021/05/09 09:54

