

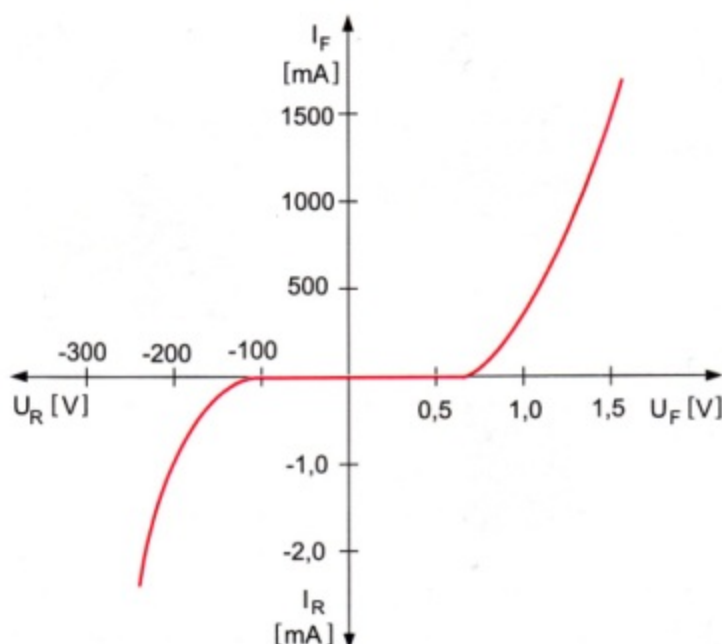
4. Podzespoły elektroniczne samochodów

4.1. Diody prostownicze

Właściwości półprzewodników niesamoistnych opisano w podrozdziale 1.2 niniejszego podręcznika. Z treści tego podrozdziału można się dowiedzieć, że w wyniku domieszkowania można uzyskać dwa rodzaje półprzewodników krystalicznych: półprzewodniki typu N oraz półprzewodniki typu P. W półprzewodnikach typu N nośnikami energii elektrycznej są elektrony, w półprzewodnikach zaś typu P energia elektryczna jest transportowana za pośrednictwem tzw. dziur. Dioda półprzewodnikowa powstanie, jeśli w trakcie procesu produkcji w półprzewodniku krystalicznym, np. w krzemie, zostanie utworzone złącze między obszarem typu P oraz obszarem typu N. Działanie diody określają właściwości złącza tych obszarów, nazywanego **złączem PN**. Złącze PN tworzą elektrony warstwy typu N przechodzące do obszaru typu P oraz dziury obszaru typu P przechodzące do obszaru typu N w chwili powstawania złącza. Energia nośników większościowych, tzn. elektronów obszaru typu N oraz dziur obszaru typu P przenikających przez złącze, musi być większa od wartości minimalnej określającej tzw. energetyczne pasmo przewodnictwa.

Z fizycznego punktu widzenia złącze PN jest utworzone z obszarów warstw typu P i N przylegających do tego złącza i nie stanowi odrębnej struktury polikrystalicznej. Złącze PN tworzą ładunki zgromadzone w obszarze przyłączeniowym i powodujące ujemną polaryzację warstwy typu P oraz dodatnią polaryzację obszaru typu N. Ładunki elektryczne tworzące złącze PN stanowią barierę potencjałów przeciwdziałającą migracji ładunków dodatnich z obszaru P do N oraz ładunków ujemnych przemieszczających się w kierunku przeciwnym. Wpływ bariery potencjałów złącza PN na ruch ładunków między obszarami tworzącymi złącze można odwzorować za pomocą siły elektromotorycznej E_j , której bieguny dodatni i ujemny są dołączone odpowiednio do przyłączeniowych warstw obszarów typu N i P (rys. 4.1a). Prąd płynący przez diodę spolaryzowaną w kierunku przewodzenia jest nazywany **prądem przewodzenia** I_F .

Aby spowodować przepływ prądu przez złącze PN, należy do diody dołączyć zewnętrzne źródło napięcia $E > E_j$ o polaryzacji przeciwnej do polaryzacji bariery potencjałów. W tym celu dodatni (+) biegun tego źródła musi być dołączony



Rys. 4.3. Przykład przebiegu charakterystyki prądowo-napięciowej diody prostowniczej

dzenia $R_F = \frac{U_F}{I_F} = \frac{2U_P}{I_F} = \frac{2 \cdot 0,8}{1000 \cdot 10^{-3}} = 1,6 \, \Omega$. Po spolaryzowaniu diody w kierunku zaporowym napięciem $U_R = 200 \, \text{V}$, prąd wsteczny płynący przez diodę ma wartość ok. $I_R = 1 \, \text{mA}$ i dlatego rezystancja diody zwiększa się do wartości $R_R = \frac{U_R}{I_R} = \frac{200}{1 \cdot 10^{-3}} = 200\,000 \, \Omega = 200 \, \text{k}\Omega$. Ze względu na duże zróżnicowanie wartości rezystancji diody spolaryzowanej w kierunku przewodzenia oraz w kierunku zaporowym, diody wykorzystuje się w układach prostujących, zamieniających prąd przemienny (bipolarny) na prąd jednokierunkowy (unipolarny).

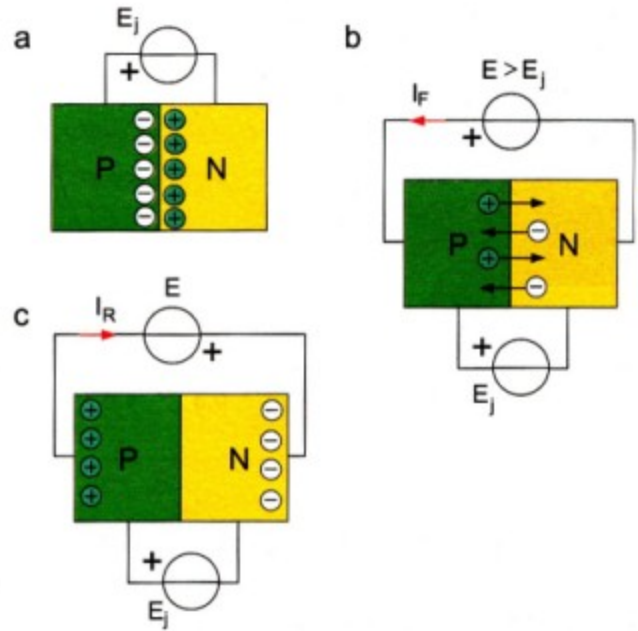
Mała wartość rezystancji diody spolaryzowanej w kierunku przewodzenia powoduje, że przepływowi prądu I_F przez diodę towarzyszy wydzielanie na złączu PN mocy cieplnej o wartości $P = U_F I_F$. Warunkiem prawidłowej pracy diody jest dobre odprowadzanie ciepła z tego złącza. W przypadku niespełnienia tego warunku, przy braku odpowiedniego rozpraszania wydzielanego ciepła wzrasta temperatura złącza i są niszczone wiązania walencyjne atomów w krystalicznej strukturze diody. Wzrost wartości prądu płynącego przez diodę ponad dopuszczalną wartość powoduje zwiększenie temperatury złącza, czego skutkiem może być **przegrzanie złącza** lub **przebiecie termiczne**. Termiczne uszkodzenie złącza może wystąpić przy polaryzacji zarówno w kierunku przewodzenia, jak i w kierunku zaporowym.

Złącza PN współczesnych diod półprzewodnikowych wytwarza się w monokryształach krzemu. Dopuszczalna temperatura pracy diod krzemowych nie przekracza 220°C . Diody te mają złącze o szerokości nieprzekraczającej $10 \, \mu\text{m}$. Przy takiej szerokości złącza napięcie zaporowe o wartości $200 \, \text{V}$ jest przyczyną występowania na złączu pola elektrycznego o natężeniu $200 \, \text{MV/m}$, które po-

do anody, biegun zaś ujemny (–) do katody diody (rys. 4.1b). Taka polaryzacja elektrod diody nazywa się **polaryzacją w kierunku przewodzenia**.

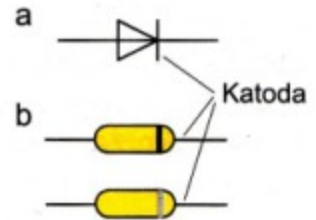
Dołączenie do diody źródła napięcia o kierunku zgodnym z polaryzacją bariery potencjałów nazywa się **polaryzacją w kierunku zaporowym** (rys. 4.1c). Przy takiej polaryzacji przez diodę przepływa prąd o małej wartości zwiększający szerokość bariery potencjałów, który jest nazywany **prądem wstecznym** lub **zaporowym**.

Symbol graficzny diody stanowi strzałka (rys. 4.2a), określająca kierunek przepływu prądu przewodzenia, oraz kreska, oznaczająca katodę. W celu ułatwienia polaryzacji w kierunku przewodzenia, katodę diod rzeczywistych oznacza się na obudowie literą *K* lub pojedynczym, srebrnym albo czarnym paskiem, umieszczonym od strony wyprowadzenia katody (rys. 4.2b).



Rys. 4.1. Złącze PN niespolaryzowane (a), spolaryzowane w kierunku przewodzenia (b) oraz spolaryzowane w kierunku zaporowym (c)
 I_F – prąd przewodzenia, I_R – prąd wsteczny

Rys. 4.2. Oznaczenia diody półprzewodnikowej
 a – symbol graficzny, b – oznaczenia kierunku przewodzenia

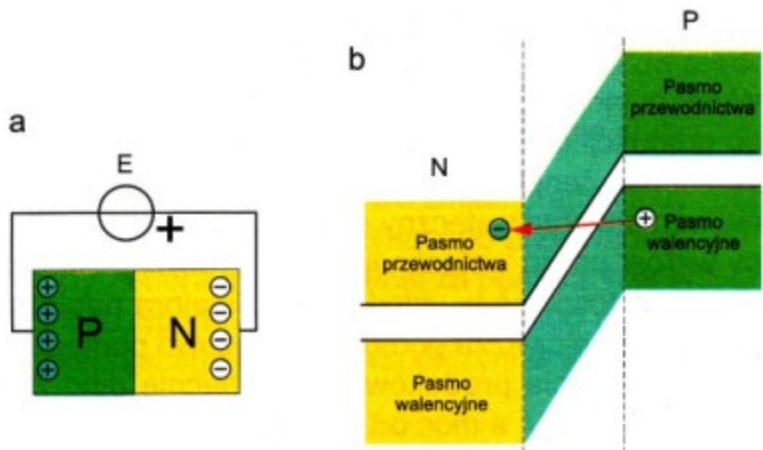


Zależność prądu płynącego przez diodę od napięcia dołączonego do jej elektrod jest nieliniowa, bez względu na kierunek tego napięcia. Przebieg tej charakterystyki dla typowej diody krzemowej przedstawiono na rys. 4.3.

Z charakterystyki przedstawionej na rys. 4.3 wynika, że pokonanie bariery potencjałów złącza PN wymaga dołączenia do elektrod diody napięcia o małej wartości, które jest nazywane **napięciem progowym** U_p . Wartość napięcia progowego diod germanowych wynosi od 0,15 do 0,30 V, a diod krzemowych od 0,6 do 0,8 V.

Z charakterystyki prądowo-napięciowej $I_F = f(U_F)$ (rys. 4.3) diody spolaryzowanej w kierunku przewodzenia wynika, że dołączenie do elektrod diody napięcia U_F o wartości dwukrotnie większej od napięcia progowego powoduje przepływ prądu I_F o wartości ok. 1000 mA. Z przedstawionych wartości napięcia U_F i prądu I_F wynika, że rezystancja diody spolaryzowanej w kierunku przewo-

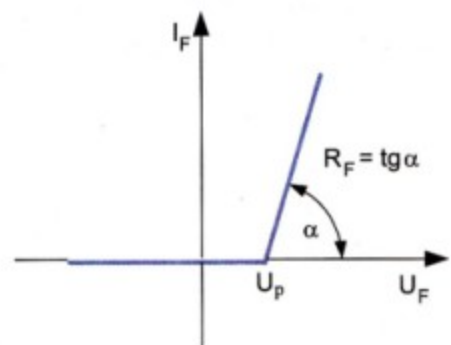
Rys. 4.4. Złącze PN spolaryzowane zaporowo (a) i poziomy warstw energetycznych diody podczas przebicia złącza (b)



woduje zmiany energetyczne w obszarze przyłączowym, przedstawione na rys. 4.4.

Przy takiej konfiguracji poziomów energetycznych elektrony walencyjne z obszaru typu N przechodzą przez barierę potencjałów złącza do obszaru typu P i stają się elektronami swobodnymi o energii należącej do pasma przewodnictwa tego obszaru. Wartość napięcia, przy którym następuje taka rekonfiguracja energetyczna złącza, nazywa się **napięciem przebicia**. Zjawisko przepływu skrótnego elektronów przez złącze w zwykłych diodach może doprowadzić do uszkodzenia bariery potencjałów. Przyczyną tego uszkodzenia jest przepływ **prądu lawinowego**. Prąd ten powstaje pod wpływem pola elektrycznego o dużym natężeniu. Pod wpływem tego pola elektrony pochodzące z obszaru typu N uderzają w atomy struktury krystalicznej obszaru typu P i wytwarzają dodatkowe pary elektron-dziura, których liczba wzrasta lawinowo. Gwałtowny (lawinowy) przyrost liczby nośników ładunku elektrycznego jest przyczyną praktycznie niekontrolowanego zwiększenia prądu płynącego przez złącze. Wzrostowi wartości skutecznej prądu lawinowego towarzyszy wzrost temperatury monokryształu diody niszczącej złącze PN.

Charakterystyka prądowo-napięciowa diody półprzewodnikowej jest nieliniowa. Do celów praktycznych może być jednak aproksymowana odcinkami linii prostej. Po uwzględnieniu napięcia bariery potencjałów U_p oraz rezystancji R_F w kierunku przewodzenia, przebieg aproksymowanej odcinkami linii prostej charakterystyki prądowo-napięciowej diody przedstawiono na rys. 4.5.

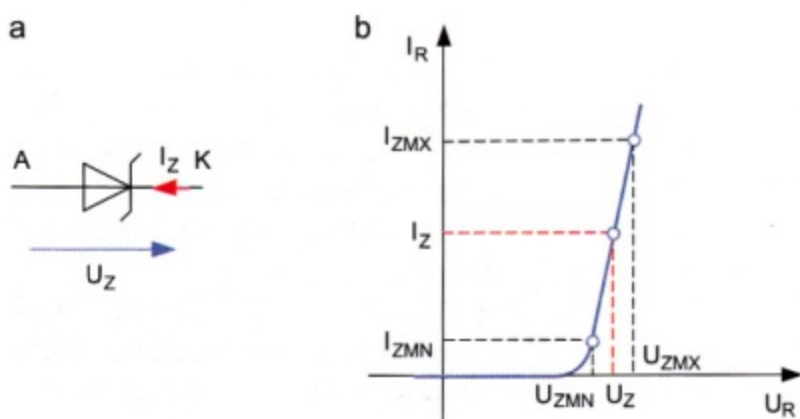


Rys. 4.5. Aproksymowana odcinkami linii prostej charakterystyka prądowo-napięciowa diody

4.2. Diody Zenera

Występujące przy polaryzacji wstecznej przebiecie złącza PN wykorzystano do produkcji diod stabilizujących napięcie, które są nazywane diodami Zenera. W diodach Zenera wzrost napięcia zaporowego do wartości, przy której gwałtownie wzrasta prąd wsteczny, wykorzystano do stabilizacji napięcia. Niekontrolowany wzrost prądu wstecznego może również zniszczyć diodę Zenera. Aby taką diodę zastosować w układzie stabilizatora napięcia, prąd płynący przez diodę należy ograniczyć za pomocą szeregowo podłączonego rezystora. Wartość napięcia przebicia produkowanych obecnie diod Zenera mieści się w zakresie od 2 V do 200 V, a moc od 0,25 W do 50 W.

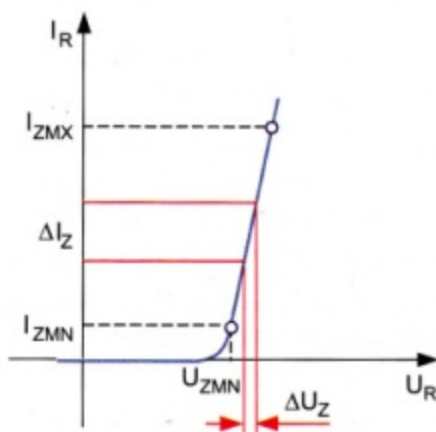
Oznaczenie diody Zenera oraz przebieg charakterystyki prądowo-napięciowej takiej diody spolaryzowanej zaporowo przedstawiono na rys. 4.6.



Rys. 4.6. Dioda Zenera
a – symbol graficzny,
b – charakterystyka
prądowo-napięciowa

Diodę Zenera najczęściej wykorzystuje się jako wzorzec napięcia stałego lub stabilizator tego napięcia. Wówczas przez diodę płynie prąd I_R o wartości większej lub równej I_{ZMN} oraz mniejszej od I_{ZMX} . Pierwsza z wymienionych wartości prądu wynika z zakrzywienia charakterystyki $I_R = f(U_R)$, drugą zaś określa maksymalna moc wydzielana na diodzie. W wymienionym zakresie zmian prądu

wstecznego charakterystykę prądowo-napięciową diody można zastąpić odcinkiem linii prostej o nachyleniu określonym rezystancją dynamiczną $r_z = \frac{\Delta U_Z}{\Delta I_Z}$, osiągającą wartość



Rys. 4.7. Wyznaczanie rezystancji dynamicznej diody Zenera

około kilku omów. W praktyce oznacza to, że dużym zmianom wartości prądu ΔI_Z towarzyszą małe zmiany ΔU_Z napięcia na diodzie. Dzięki temu dioda Zenera może pełnić rolę stabilizatora napięcia. W zakresie stabilizacji napięcia diodę taką można przedstawić jako szeregowo połączenie idealnego źródła napięcia U_Z oraz opornika o rezystancji równej r_z (rys. 4.7).

W odróżnieniu od rezystancji dynamicznej, rezystancję statyczną R_z w danym punkcie pracy diody Zenera wyznacza stosunek napięcia na diodzie do prądu płynącego przy tym napięciu przez diodę, tzn.: $R_z = \frac{U_z}{I_z}$. W zakresie stabilizacji

napięcie na diodzie pozostaje stałe w szerokim zakresie zmian prądu diody i dla tego rezystancję statyczną określa wartość prądu płynącego przez diodę.

Właściwości stabilizujące diody Zenera można ocenić za pomocą tzw. współczynnika stabilizacji, określającego wpływ zmian prądu płynącego przez diodę na zmiany spadku napięcia występującego na niej. Współczynnik stabilizacji definiuje zależność:

$$K_z = \frac{\frac{\Delta I_z}{I_z}}{\frac{\Delta U_z}{U_z}} = \frac{R_z}{r_z} \quad (4.1)$$

Z zależności (4.1) wynika, że wartość współczynnika stabilizacji jest określona stosunkiem rezystancji statycznej do rezystancji dynamicznej. Większa wartość tego współczynnika oznacza lepszą przydatność diody do stabilizacji napięcia.

W technice samochodowej diody Zenera wykorzystuje się w układach stabilizujących napięcie zasilające podzespoły elektroniczne, jako diody prostujące napięcie wyjściowe alternatorów oraz w przekaźnikach ochrony przeciwprzepięciowej. Zadaniem diod stosowanych w tego rodzaju przekaźnikach jest ochrona podzespołów elektronicznych przed skutkami impulsowych przepięć napięciowych, występujących podczas przełączania elektromechanicznych podzespołów stykowych, tj. przekaźników oraz kontaktronów. W układach ochrony przeciwprzepięciowej należy stosować zabezpieczające diody Zenera, charakteryzujące się dużą odpornością na prądy wsteczne o dużej wartości szczytowej.

4.3. Diody pojemnościowe

Ładunek przestrzenny występujący po obu stronach złącza PN rozdziela bariera potencjałów charakteryzująca się przyzłaczowym rozkładem przestrzennym o pewnej szerokości d . Z tego względu bariera ta zachowuje się podobnie jak kondensator płaski o pojemności kilku pikofaradów, której wartość określa zależność:

$$C_{PN} = \frac{\varepsilon S}{d} \quad (4.2)$$

W zależności tej ε oraz S oznaczają odpowiednio bezwzględną przenikalność elektryczną materiału półprzewodnika oraz powierzchnię złącza.

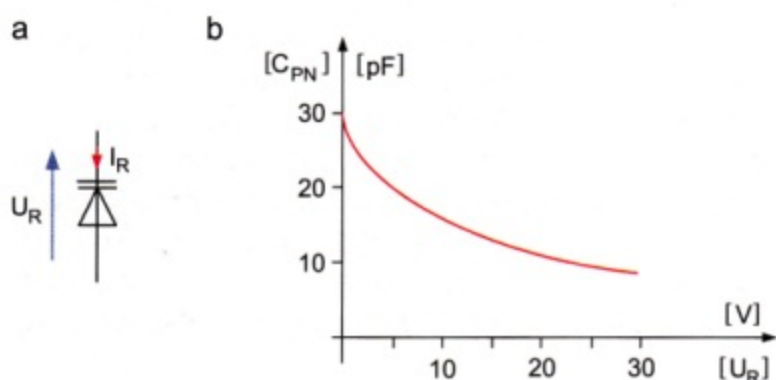
Pojemność złącza PN maleje wraz ze wzrostem napięcia wstecznego na tym złączu, ponieważ zwiększa się szerokość bariery potencjałów. Między pojemnością złącza a napięciem U_R polaryzującym wstecznie złącze występuje zależność:

$$C_{PN} = \frac{C_0}{\left(1 + \frac{U_R}{\varphi_T}\right)^{\frac{1}{n}}} \quad (4.3)$$

W zależności tej C_0 oznacza pojemność złącza przy braku napięcia polaryzującego wstecznie złącze PN, n określa parametry technologiczne złącza zależne od rozkładu koncentracji domieszek i może przyjmować wartości od 1 do 3,

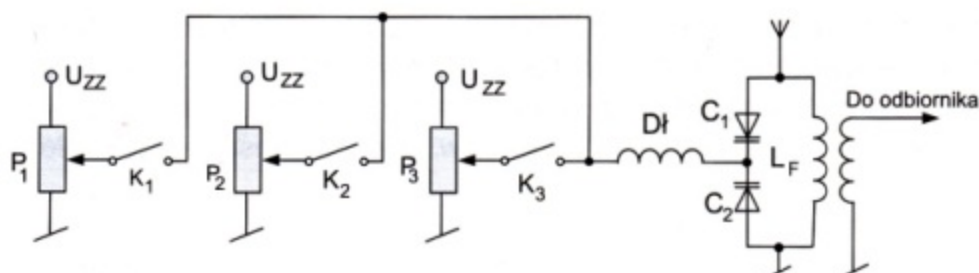
$\varphi_T = \frac{kT}{q}$ jest potencjałem elektrokinetycznym (k jest stałą Boltzmana, T oznacza temperaturę złącza, q – ładunek elektronu).

Diody półprzewodnikowe o takich właściwościach nazywa się diodami pojemnościowymi lub warikapami o sterowanej napięciem wstecznym pojemności C_{PN} . Symbol graficzny takiej diody oraz charakterystykę $C_{PN} = f(U_R)$ przedstawiono na rys. 4.8.



Rys. 4.8. Symbol diody pojemnościowej (a) oraz jej charakterystyka (b)

Diody pojemnościowe w technice samochodowej wykorzystuje się w układach o dużej częstotliwości, np. nadajnikach i odbiornikach sygnału radiowego, za pomocą którego jest uruchamiany centralny zamek. W tego rodzaju układach dio-



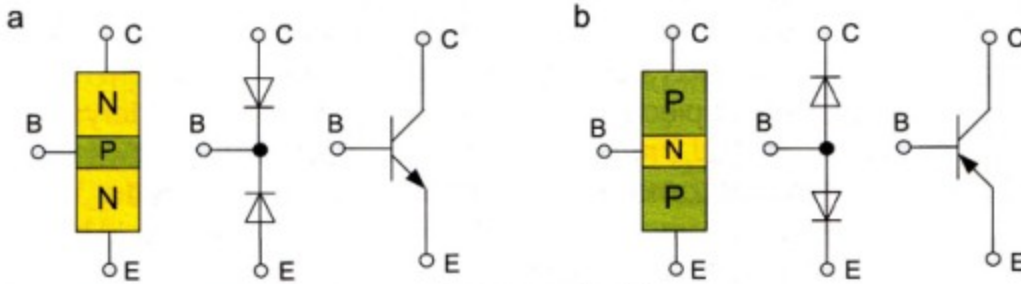
Rys. 4.9. Programator stacji radiowych w zakresie UKF, wykorzystujący diody pojemnościowe

dy pojemnościowe dostrajają obwody rezonansowe wzmacniaczy parametrycznych do częstotliwości fali radiowej. Obwód rezonansowy odbiornika takich sygnałów, z diodami pojemnościowym pełniącymi rolę kondensatorów o dokładnie dobranej pojemności, przedstawiono na rys. 4.9. Wartość pojemności warikapów określa napięcia występujące na suwaku potencjometru P.

4.4. Tranzystory bipolarne

Tranzystor bipolarny składa się z trzech warstw półprzewodnika rozdzielonych dwoma złączami PN. Warstwy zewnętrzne nazywa się emiterem (E) i kolektorem (C), a warstwa wewnętrzna nosi nazwę bazy (B) – rys. 4.10. Jeśli baza tranzystora jest półprzewodnikiem typu P, a warstwy zewnętrzne są półprzewodnikami typu N, to tranzystor jest typu N-P-N. W odwrotnym przypadku, tzn. gdy bazę stanowi półprzewodnik typu N, tranzystor jest typu P-N-P.

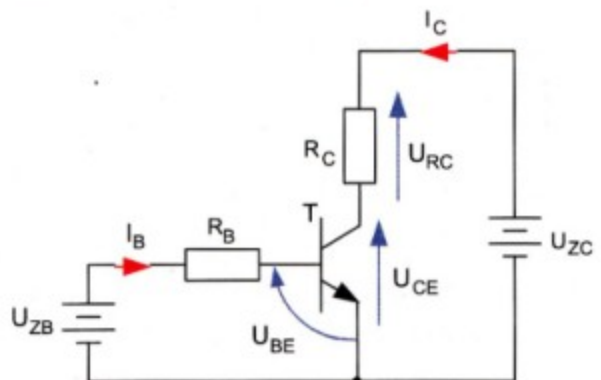
Wielu autorów do wyjaśnienia działania tranzystorów bipolarnych, oznaczanych w literaturze angielskojęzycznej symbolem BJT, wykorzystuje ich model, który stanowią dwie przeciwstawnie połączone diody (rys. 4.10).



Rys. 4.10. Tranzystory bipolarne N-P-N (a) i P-N-P (b)

Warunkiem przewodzenia prądu przez tranzystor bipolarny jest polaryzacja złącza baza-emiter w kierunku przewodzenia, a złącza baza-kolektor w kierunku zaporowym. Przykład polaryzacji tranzystora N-P-N za pomocą źródeł napięcia podłączonych do emitera tego tranzystora przedstawiono na rys. 4.11.

Przy polaryzacji tranzystora w sposób przedstawiony na rys. 4.11, przez złącze kolektor-baza przepływa prąd mniejszościowy, określony wartością napięcia polaryzującego złącze baza-emiter w kierunku przewodzenia. Jeśli szerokość bazy jest mała, małe zmiany wartości tego napięcia powodują duże zmiany prądu złącza kolektorowego spolaryzowa-



Rys. 4.11. Polaryzacja tranzystora za pomocą źródeł napięcia podłączonych do emitera

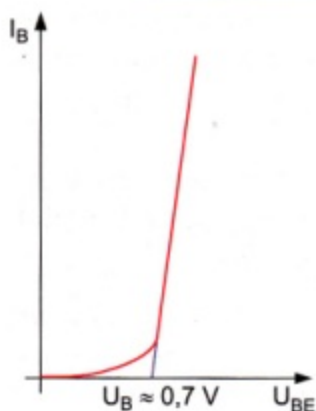
nego zaporowo. Niewielkie zmiany napięcia występującego na złączu baza-emiter są przyczyną niewielkich zmian prądu płynącego przez to złącze i dlatego w tranzystorze BJT występuje efekt wzmacniania prądu, którego miarą jest statyczny współczynnik wzmacnienia

$$\beta = \frac{I_C}{I_B} \quad (4.4)$$

W zależności tej I_C oraz I_B oznaczają odpowiednio prąd kolektora oraz prąd bazy tranzystora. Współczynnik β przyjmuje wartości od kilkudziesięciu do kilkuset.

W układzie przedstawionym na rys. 4.11 prąd emitera I_E określa zależność:

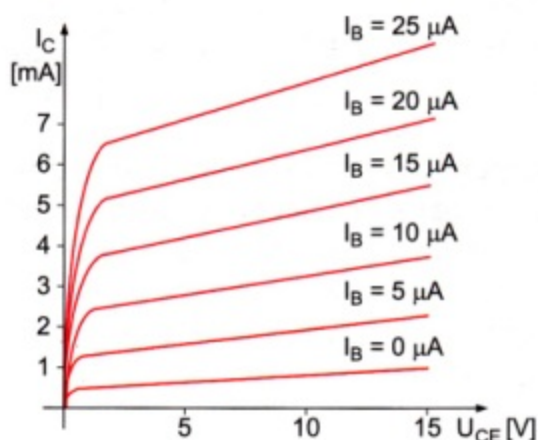
$$I_E = I_B + I_C \quad (4.5)$$



Rys. 4.12. Charakterystyka wejściowa tranzystora BJT

Od strony zacisków B-E tranzystor jest widziany jako złącze PN spolaryzowane w kierunku przewodzenia, o charakterystyce prądowo-napięciowej $I_B = f(U_{BE})$ mającej przebieg zgodny z charakterystyką diody półprzewodnikowej (rys. 4.12).

W tranzystorach BJT wartość prądu kolektora jest regulowana za pomocą napięcia U_{BE} i dlatego charakterystyka $I_B = f(U_{BE})$ jest nazywana charakterystyką wejściową tych tranzystorów. Obwodem wyjściowym tranzystora jest rezystor włączony między jego kolektorem a źródłem napięcia zasilającego U_{ZC} . Dlatego charakterystyka $I_C = f(U_{CE})$ nazywa się charakterystyką wyjściową i jest wyznaczana dla ustalonych wartości prądu bazy (rys. 4.13) lub ustalonego napięcia baza-emiter.

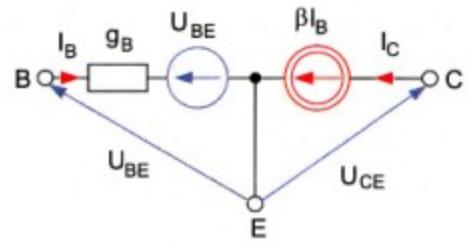


Rys. 4.13. Charakterystyki wyjściowe tranzystora BJT

Z przebiegu charakterystyk przedstawionych na rys. 4.13 wynika, że wraz ze wzrostem od zera napięcia U_{CE} wartość prądu kolektora I_C szybko rośnie, gdy $I_B > 0$, ponieważ oba złącza PN tranzystora są spolaryzowane w kierunku przewodzenia. Wartość prądu I_C stabilizuje się praktycznie na stałym poziomie równym βI_B , gdy napięcie kolektor-baza maleje do zera, tzn. dla $U_{CB} = 0$. Taki stan pracy tranzystora jest nazywany stanem aktywnym, z racji wzmacniania prądu

bazy. W tym stanie, dla ustalonego punktu pracy, właściwości tranzystora N-P-N można odwzorować za pomocą schematu zastępczego, przedstawionego na rys. 4.14.

Na schemacie tym elementy g_B oraz $U_{BE} = 0,7 \text{ V}$ odwzorowują charakterystykę prądowo-napięciową spolaryzowanego w kierunku przewodzenia złącza ba-



Rys. 4.14. Schemat zastępczy tranzystora N-P-N

za-emiter tranzystora N-P-N. Parametr $g_B = \frac{I_B}{\varphi_T}$ jest konduktancją bazy, gdzie

$\varphi_T = \frac{kT}{q}$ jest potencjałem elektrokinetycznym (k jest stałą Boltzmana, T oznacza temperaturę złącza, q – ładunek elektronu). W aktywnym stanie pracy tranzystor BJT jest sterowanym źródłem prądowym, stąd na schemacie zastępczym (rys. 4.14) występuje źródło prądowe o wydajności βI_B .

Schemat zastępczy na rys. 4.14 odwzorowuje również właściwości tranzystorów P-N-P. Dla takich tranzystorów źródło napięciowe U_{BE} oraz kolektorowe źródło prądowe mają kierunek odwrotny, w porównaniu z kierunkiem zaznaczonym na rys. 4.14. Gdy oba złącza PN są spolaryzowane w kierunku przewodzenia, tranzystor jest nasycony i prąd bazy nie ma wpływu na wartość prądu kolektora. Prąd ten jest określony przez napięcie zasilające U_{CC} oraz rezystancję opornika R_C stanowiącego obciążenie kolektora tranzystora: $I_C = \frac{U_{CC}}{R_C}$. Przy

wyznaczaniu prądu płynącego przez nasycony tranzystor można pominąć napięcie U_{CE} , ponieważ ma ono wartość bliską zera, oznaczaną w katalogach symbolem U_{Csat} . Tranzystor BJT będzie nasycony dla prądu bazy I_B o wartości spełniającej nierówność $\beta I_B > I_C$. W tym stanie tranzystory takie zachowują się jak źródła napięciowe, na których występuje napięcie nasycenia U_{Csat} . Przy zerowej wartości prądu bazy ($I_B = 0$) tranzystor jest odcięty i prąd kolektora ma wartość prądu zerowego I_{C0} . Tak więc oba złącza PN tranzystora nie przewodzą prądu i dlatego napięcie kolektor-emiter ma wartość zbliżoną do napięcia zasilającego: $U_{CE} \cong U_{CC}$.

4.5. Tranzystory polowe

Tranzystory polowe (FET – ang. *Field Effect Transistors*) stanowią klasę przyrządów półprzewodnikowych, w których prąd płynie wzdłuż wąskiego kanału typu P lub N, od źródła S do drenu D. Szerokość kanału przewodzącego prąd jest regulowana polem elektrycznym o natężeniu określonym wartością napięcia dołączonego do elektrody sterującej, zwanej bramką G. Stąd wywodzi się nazwa tych tranzystorów. Znane są dwie grupy tranzystorów polowych. Do pierwszej grupy należą tranzystory polowe złączowe, których bramkę stanowi spolaryzo-

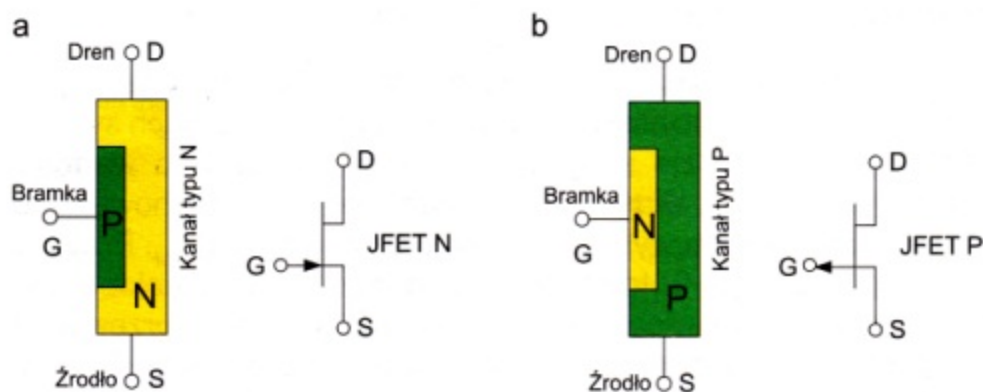
wane zaporowo złącze PN wnikające do kanału. Drugą grupą są tranzystory z bramką odizolowaną od kanału cienką warstwą tlenku krzemu, znane pod nazwą MOSFET (ang. *Metal Oxide Semiconductor FET*). W tej grupie można wyodrębnić dwie podgrupy: tranzystory D-MOSFET (ang. *Depletion-MOSFET*) o kanale wzbogaconym, który przewodzi prąd bez konieczności polaryzacji bramki, oraz tranzystory E-MOSFET (ang. *Enhancement-MOSFET*) o kanale zubożonym, przewodzące prąd pod wpływem napięcia bramki.

Z energetycznego punktu widzenia tranzystory polowe charakteryzują się wspólną pozytywną cechą, którą jest duża rezystancja wejściowa.

Tranzystory polowe, ze względu na możliwość wykonywania ich w dużej skali integracji, są podzespołami współczesnych systemów mikroprocesorowych produkowanych jako autonomiczne układy scalone.

4.5.1. Tranzystory złączowe (JFET)

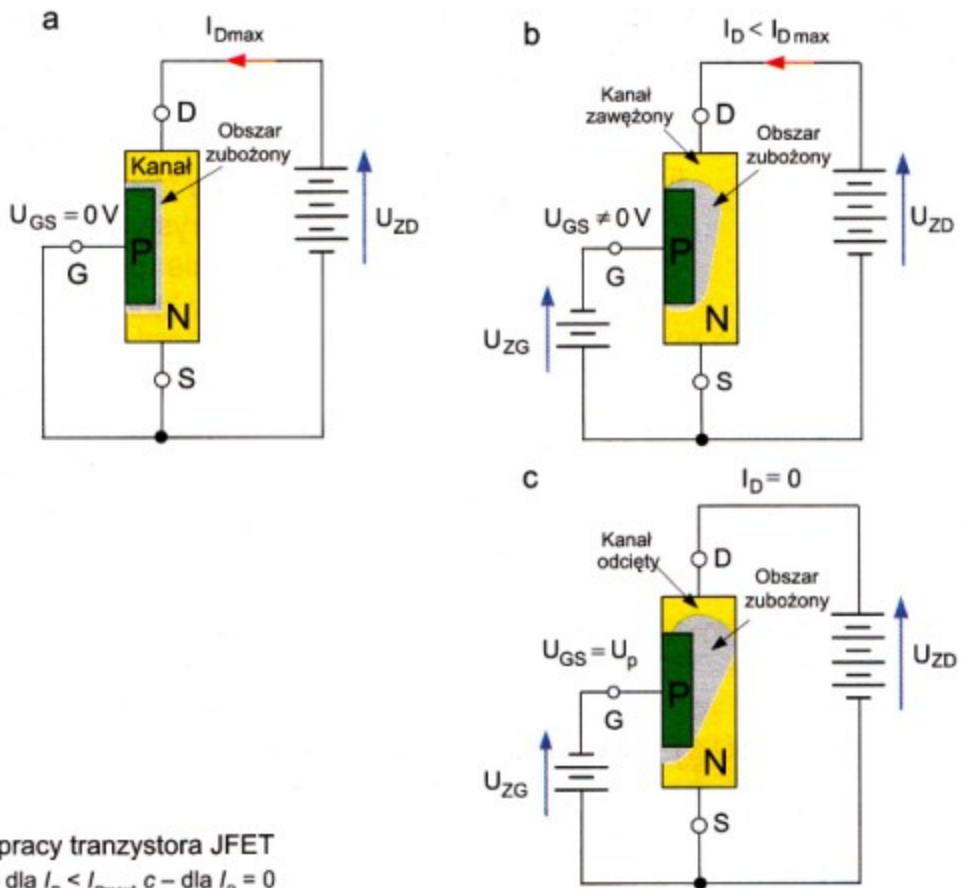
Konstrukcję tranzystorów polowych złączowych wyjaśniono na rys. 4.15. Kanał tych tranzystorów może być wykonany z półprzewodnika typu P lub typu N. Końce kanału stanowią dwie napylone metalowe elektrody: źródło S oraz dren D. W celu utworzenia złącza PN bramki, w kanał jest dyfundowany obszar półprzewodnika o przeciwnym rodzaju nośników, na który jest napylona metalowa elektroda bramki G.



Rys. 4.15. Struktura oraz symbole tranzystorów polowych złączowych z kanałem typu N (a) oraz z kanałem typu P (b)

W celu wymuszenia przepływu prądu przez kanał między źródłem a drenem należy dołączyć napięcie wymuszające przepływ nośników większościowych. Jeśli tranzystor ma kanał typu N, to dren musi mieć potencjał dodatni względem źródła (rys. 4.16). Dren tranzystorów z kanałem typu P ma polaryzację ujemną.

Przy zerowym napięciu bramki ($U_{GS} = 0$) prąd kanału I_D ma wartość maksymalną, ponieważ pozbawiony nośników większościowych obszar w pobliżu złącza PN bramki zajmuje objętość kanału najmniejszą z możliwych (rys. 4.16a). Zwiększanie objętości tego obszaru za pomocą napięcia bramki powoduje zmniejszanie przekroju kanału, czego skutkiem jest zmniejszanie wartości prądu płynącego między źródłem a drenem tranzystora (rys. 4.16b). Zwiększenie na-

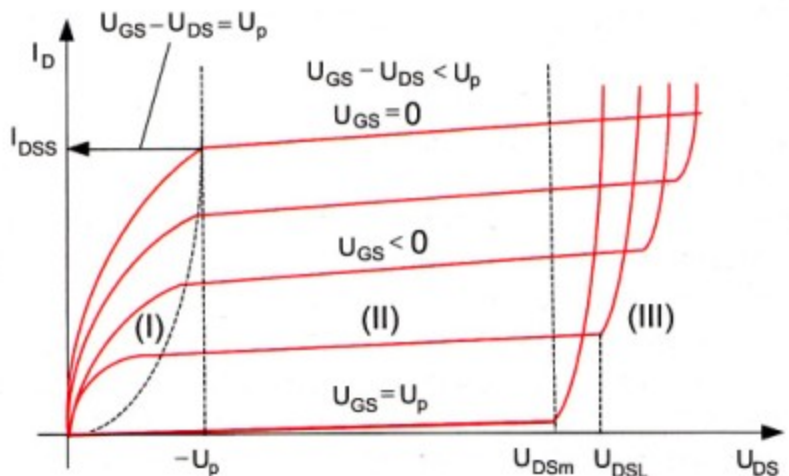


Rys. 4.16. Stany pracy tranzystora JFET

a – dla $I_D = I_{Dmax}$, b – dla $I_D < I_{Dmax}$, c – dla $I_D = 0$

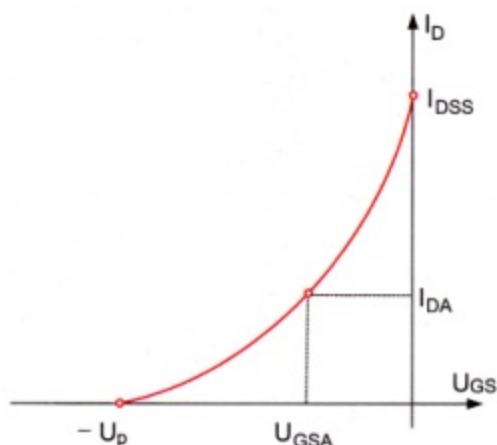
pięcia U_{GS} do wartości napięcia odcięcia U_p powoduje, że cały przekrój kanału zajmuje obszar zubożony i prąd drenu I_{DS} przestaje płynąć (rys. 4.16c). Przedstawione zjawiska mają wpływ na przebieg charakterystyk wyjściowych $I_D = f(U_{DS})$ tranzystora JFET (rys. 4.17).

W przebiegu charakterystyk wyjściowych można wyróżnić trzy obszary. Pierwszy obszar, w którym prąd drenu I_D zmienia wartość praktycznie liniowo wraz ze wzrostem napięcia U_{DS} , nazywa się obszarem pracy omowej lub triodowej. W tym obszarze pracy prąd drenu jest sterowany napięciem bramki i powoduje



Rys. 4.17. Charakterystyki wyjściowe tranzystora JFET typu N

dodatkową polaryzację wsteczną złącza bramka-kanal napięciem $U_{GS} - U_{DS} < U_p$. Zaznaczona na rys. 4.17 linią przerywaną granica obszaru trydowego występuje dla $U_{GS} - U_{DS} = U_p$.



Rys. 4.18. Charakterystyka przejściowa tranzystora JFET

Napięcie bramki ma również wpływ na prąd drenu, gdy tranzystor JFET jest nasycony i charakterystyki $I_D = f(U_{DS})$ mają przebieg praktycznie poziomy. Jest to drugi obszar pracy, w którym $U_{GS} - U_{DS} > U_p$.

Trzeci obszar pracy występuje, gdy napięcie dren-źródło przekroczy wartość U_{DSL} . Po osiągnięciu tego napięcia następuje przebicie złącza bramka-kanal w pobliżu drenu, czego skutkiem jest lawinowe powielanie nośników prądu.

Oprócz charakterystyk $I_D = f(U_{DS})$, producenci tranzystorów JFET podają w katalogach tzw. charakterystykę przejściową $I_D = f(U_{GS})$ – rys. 4.18. Z przebiegu charak-

terystyki $I_D = f(U_{GS})$ można wyznaczyć, dla ustalonego punktem pracy A prądu drenu I_D tranzystora, wartość tzw. transkonduktancji $g_m = \frac{\Delta I_{DA}}{\Delta U_{DSA}}$, charakteryzu-

jącej właściwości wzmacniające tranzystora JFET. Jednostką tego współczynnika jest mS ($1 \text{ mS} = 1 \text{ mA/V}$).

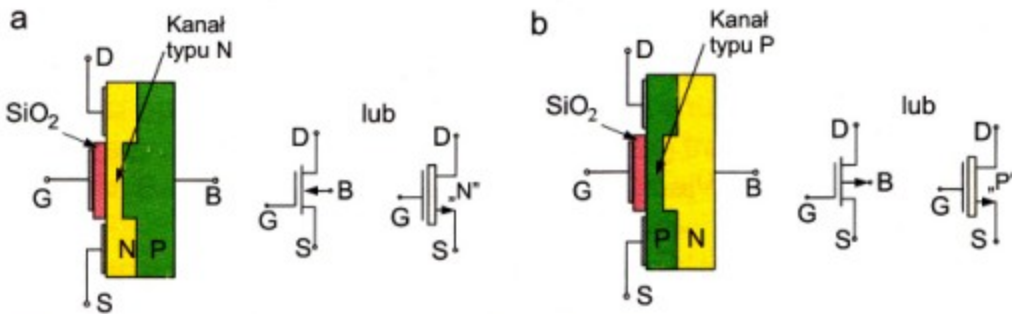
Tranzystory polowe złączowe w technice samochodowej wykorzystuje się w układach wzmacniających. Mogą one także pełnić rolę sterowanych napięciowo kluczy elektronicznych.

4.5.2. Tranzystory polowe z izolowaną bramką (MOSFET)

Oprócz tranzystorów JFET, w układach elektronicznych, w tym również w układach stosowanych w technice samochodowej, stosuje się tranzystory polowe z izolowaną bramką.

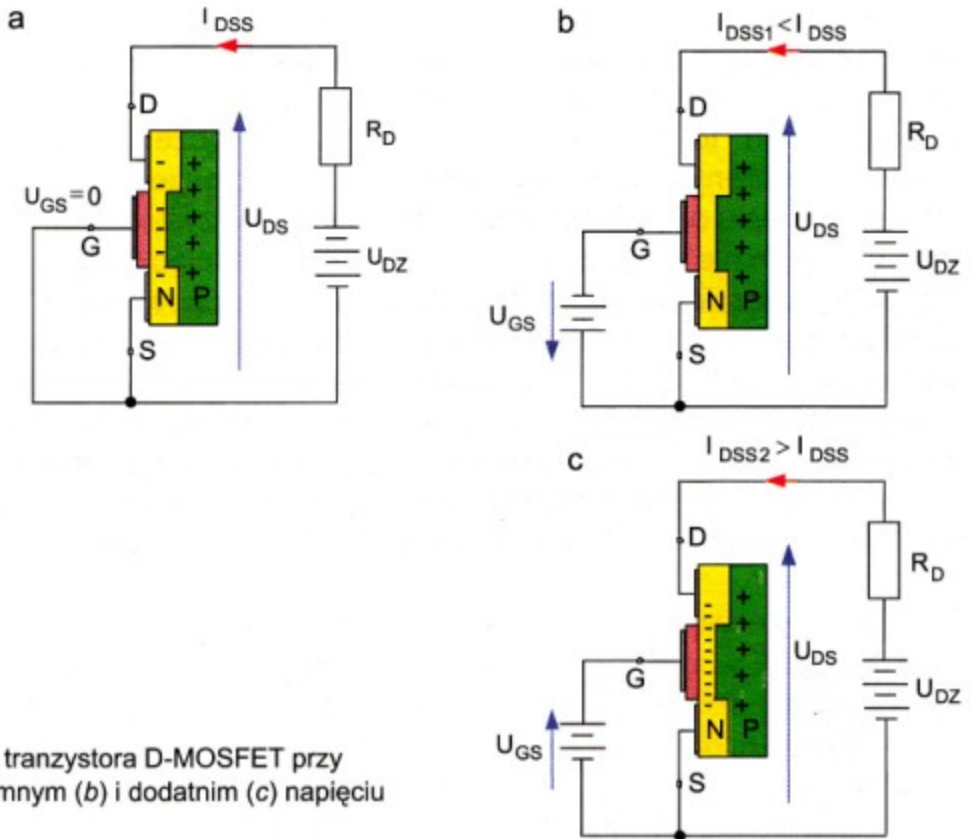
Jak wspomniano w podrozdziale 4.5, pierwszą podgrupą tranzystorów MOSFET są tranzystory D-MOSFET, z kanałem zubażonym wytworzonym na podłożu B o przeciwnym rodzaju przewodności. Budowę tych tranzystorów przedstawiono na rys. 4.19. Po obu stronach metalizowanej bramki, odizolowanej od kanału warstwą dwutlenku krzemu, bezpośrednio na kanale są naniesione dwie metalizowane elektrody: źródło S i dren D .

W tranzystorach D-MOSFET przez kanał, między źródłem a drenem, płynie prąd przy zerowym napięciu bramki. Prąd ten w katalogach jest oznaczany symbolem I_{DSS} . Ze względu na przepływ tego prądu tranzystory typu D-MOSFET nazywa się normalnie otwartymi.



Rys. 4.19. Budowa i symbol tranzystora MOSFET z kanałem typu N (a) oraz z kanałem typu P (b)

Jeżeli do bramki tranzystora doprowadzimy napięcie o polaryzacji zgodnej z rodzajem nośników kanału, to przez cienką warstwę dielektryka zaindukują się ładunki o przeciwnym znaku, zwężające przekrój kanału, co zmniejsza wartość prądu płynącego przez ten kanał (rys. 4.20).

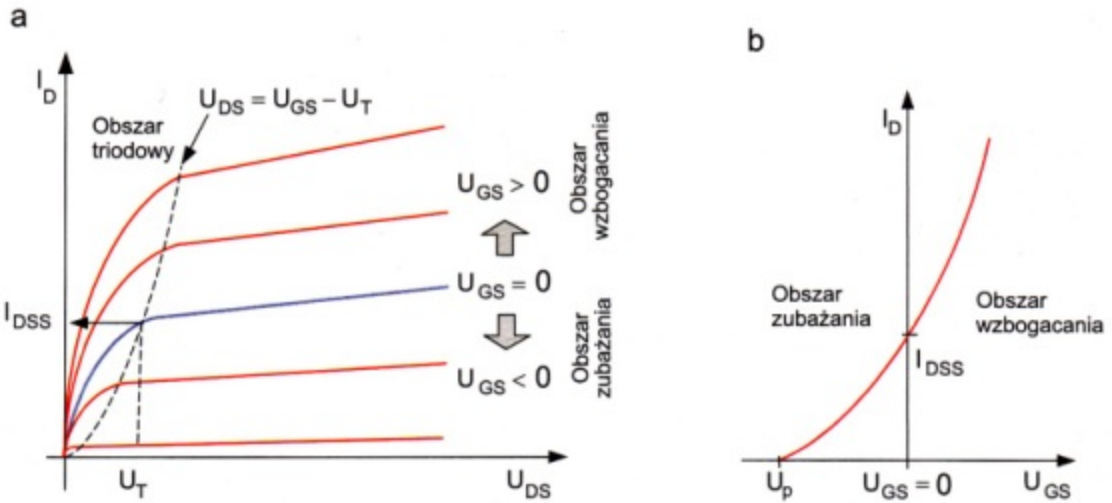


Rys. 4.20. Praca tranzystora D-MOSFET przy zerowym (a), ujemnym (b) i dodatnim (c) napięciu bramki

Po osiągnięciu przez napięcie wartości równej U_T , prąd drenu przestaje płynąć – tranzystor jest w stanie odcięcia.

Przebieg charakterystyk wyjściowych $I_D = f(U_{DS})$ oraz charakterystyki przejściowej $I_D = f(U_{GS})$ tranzystora D-MOSFET przedstawiono na rys. 4.21.

Charakterystyki wyjściowe $I_D = f(U_{DS})$ tych tranzystorów mają przebieg podobny do analogicznych charakterystyk tranzystorów JFET. W polu tych charakterystyk można również wyodrębnić obszar pracy triodowej oraz obszar nasycenia.

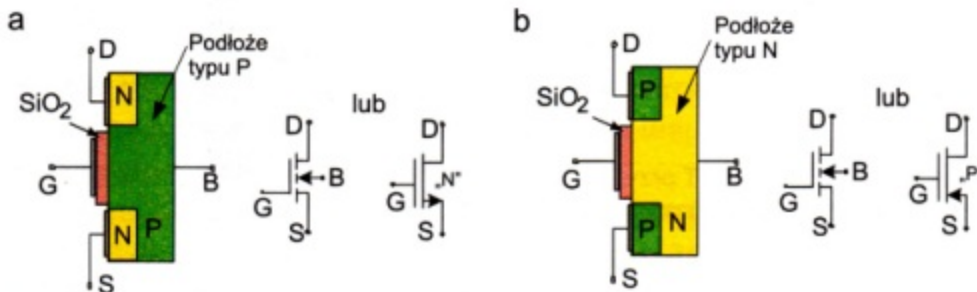


Rys. 4.21. Przebieg charakterystyki wyjściowej (a) oraz przejściowej (b) tranzystora D-MOSFET z kanałem typu N

Jedyna różnica polega na możliwości sterowania bramki napięciem zarówno dodatnim, jak i ujemnym. Jeśli napięcie na bramce ma znak zgodny ze znakiem nośników kanału, to kanał ten jest zubożany. Dołączenie do bramki napięcia o znaku przeciwnym do znaku kanału jest przyczyną jego wzbogacania.

W odróżnieniu od tranzystorów D-MOSFET, tranzystory E-MOSFET o kanale wzbogacanym nie mogą przewodzić prądu, gdy napięcie $U_{GS} = 0$, ponieważ nie mają one fizycznie utworzonego kanału. Kanał ten powstaje dopiero pod wpływem dołączenia do bramki napięcia o odpowiedniej wartości i polaryzacji. Budowę tranzystorów E-MOSFET przedstawiono na rys. 4.22.

W podłożu (np. typu P) są wdyfundowane dwie metalizowane wyspy o przeciwnym do kanału rodzaju nośników prądu elektrycznego (np. typu N): źródło S i dren D. Wyspy te rozdziela, odizolowana od podłoża warstwą dwutlenku krzemu, metalowa bramka G, umieszczona w miejscu indukowania w podłożu.



Rys. 4.22. Struktura tranzystora E-MOSFET z kanałem indukowanym typu N (a) i typu P (b)

Jeśli do bramki dołączymy napięcie o większej od U_T polaryzacji zgodnej z rodzajem nośników prądu elektrycznego podłoża, to warstwa izolacyjna SiO_2 spowoduje, że w obszarze podłoża bezpośrednio pod elektrodą bramki zostanie zaindukowany kanał przewodzący, utworzony przez nośniki o znaku zgodnym z nośnikami źródła i drenu (rys. 4.23). Wzrost napięcia bramki powoduje wzbo-

gacanie kanału i paraboliczne zwiększenie wartości prądu płynącego od źródła do drenu. Wpływ wzrostu napięcia U_{GS} na wartość prądu drenu I_D ilustruje przebieg charakterystyki przejściowej tranzystora E-MOSFET typu N, przedstawionej na rys. 4.23.

W tranzystorach E-MOSFET napięcie bramki jest przyczyną wzbogacania kanału indukowanego w podłożu. Przewodność tego kanału jest zależna od różnicy $U_{GS} - U_T$. Gdy $U_{GS} - U_T < U_{DS}$, wówczas wartość prądu drenu jest proporcjonalna do $U_{GS} - U_T$ (rys. 4.23). Jest to obszar pracy triodowej tranzystora (rys. 4.21). Po przekroczeniu przez napięcie dren-źródło granicy $U_{DS} = U_{GS} - U_T$ (rys. 4.21), prąd drenu ma wartość stałą i niezależną od wartości tego napięcia. Ten obszar pracy nazywamy obszarem nasycenia lub pentodowym.



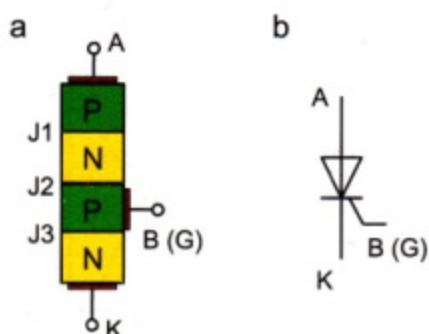
Rys. 4.23. Charakterystyka przejściowa tranzystora E-MOSFET typu N

4.6. Tyrystor i triak

W elektrycznych instalacjach samochodowych, z uwagi na wartość prądu płynącego w poszczególnych obwodach, można wyodrębnić układy i obwody, przez które płyną prądy o wartościach nieprzekraczających kilku amperów, oraz obwody i układy z prądami o wartościach kilkudziesięciu i więcej amperów. Pierwszą grupę obwodów potocznie nazywa się obwodami niskoprądowymi, a drugą grupę – obwodami silnoprądowymi. Przedstawiony podział szczególnie uwydatnia się w samochodach hybrydowych spalinowo-elektrycznych.

Do komutacji, czyli przełączania prądu w obwodach niskoprądowych, można wykorzystać kontaktrony, przekaźniki elektromagnetyczne oraz klucze elektroniczne z tranzystorami unipolarnymi lub bipolarnymi. Komutacja prądu o wartościach rzędu kilkudziesięciu amperów nie jest możliwa za pomocą wymienionych podzespołów elektromechanicznych i elektronicznych. Do tego celu można wykorzystać półprzewodnikowe elementy mocy. Do grupy tych elementów należą tyrystory i triaki. Tyrystory są łącznikami półprzewodnikowymi przystosowanymi do komutacji prądów o wartości od 0,4 do 4500 A, przy napięciu wstecznym od 50 V do 8000 V. Napięcie wsteczne określa wartość napięcia roboczego tego podzespołu, ponieważ występuje ono między katodą i anodą spolaryzowanego zaporowo tyrystora.

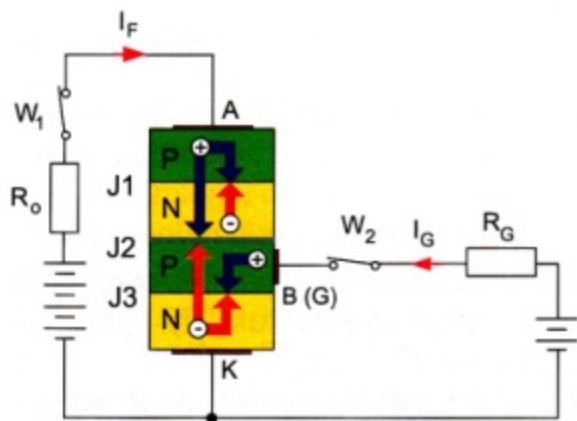
Budowę tyrystora pokazano na rys. 4.24. Typowy tyrystor składa się co najmniej z czterech warstw półprzewodnika o różnych typach przewodnictwa, uszeregowanych kolejno P-N-P-N.



Rys. 4.24. Struktura tyrystora (a) oraz jego symbol graficzny (b)

Zewnętrzne warstwy półprzewodnika P i N nazywa się odpowiednio anodą A i katodą K. Półprzewodnik wewnętrzny typu P pełni rolę bramki wyzwalającej B. Jeśli tyrystor będzie spolaryzowany w kierunku przewodzenia, tzn. na anodzie będzie występowało napięcie dodatnie względem katody, a napięcie bramki będzie miało wartość zerową, to zewnętrzne złącza J1 oraz J3 będą spolaryzowane w kierunku przewodzenia, a wewnętrzne złącze J2 będzie spolaryzowane w kierunku zaporowym. W tym stanie przez tyrystor nie płynie prąd. Tyrystor zaczyna przewodzić prąd, gdy zaczyna płynąć prąd bramki I_G , zmniejszający barierę potencjałów złącza J2.

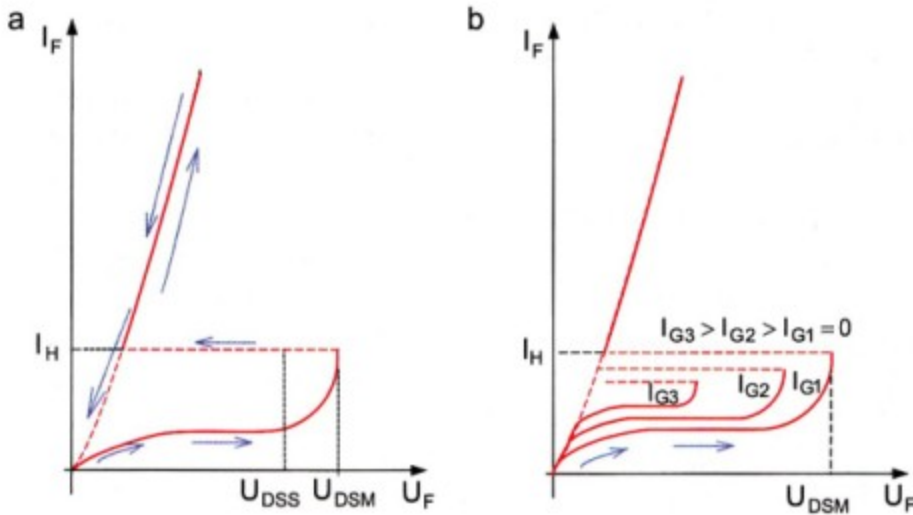
Zjawisko przewodzenia prądu przez tyrystor wynika ze zróżnicowania domieszkowania i grubości obszaru katody oraz warstwy sterującej, do której jest dołączona elektroda bramki. Koncentracja elektronów w katodzie i w mniej domieszkowanej cienkiej warstwie sterującej ($10 \dots 30 \mu\text{m}$) ma odpowiednio wartości: $10^{18}/\text{cm}^3$ i $10^{17}/\text{cm}^3$. Przepływ prądu bramki powoduje, że elektrony z warstwy katodowej dyfundują przez cienką warstwę sterującą do anody i zmniejszają rezystancję złącza J2 (rys. 4.25). Drugą przyczyną zwiększenia przewodności tego złącza są dziury pochodzące z anody, które rekombinują z elektronami warstwy zaporowej.



Rys. 4.25. Ruch nośników w strukturze tyrystora po spolaryzowaniu elektrod

Pole elektryczne wytwarzane przez źródło napięcia zasilające obwód anodowy tyrystora powoduje, że elektrony i dziury przemieszczają się między anodą i katodą tyrystora nawet po zaniku prądu bramki. Przyczyną spadku napięcia na przewodzącym tyrystorze jest rezystancja warstwy sterującej. Sposób wyzwalania tyrystora odzwierciedla przebieg jego charakterystyki prądowo-napięciowej (rys. 4.26).

Z charakterystyki przedstawionej na rys. 4.26a wynika, że w spolaryzowanym w kierunku przewodzenia tyrystorze o zerowym napięciu na bramce płynie prąd blokowania, o wartości wzrastającej wraz z napięciem polaryzującym dodatnio anodę względem katody. Gdy napięcie między anodą i katodą osiągnie wartość U_{DSM} , następuje przebicie tyrystora w kierunku przewodzenia i tyrystor przechodzi ze stanu zaporowego w stan pełnego przewodzenia. Stan taki będzie stabilny, tzn. prąd przez tyrystor będzie płynął w sposób ciągły, jeśli wartość tego prądu będzie większa od tzw. prądu podtrzymania I_H .



Rys. 4.26. Charakterystyki $I_F = f(U_F)$ tyrystora: (a) bez polaryzacji bramki, (b) dla trzech różnych prądów bramki

Tyrystory są tak zaprojektowane, aby napięcie anodowe nie przekraczało wartości $U_{DSS} \leq 0,8 U_{DSM}$.

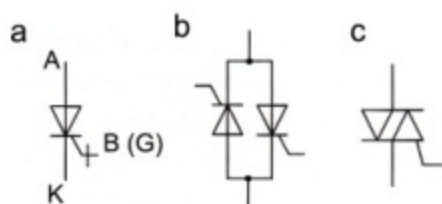
Polaryzacja bramki napięciem U_G dodatnim względem katody spowoduje, że tyrystor zacznie przewodzić prąd dla napięć anodowych U_{DS} , których wartość zmniejsza się wraz ze wzrostem wartości szczytowej impulsu prądowego I_G płynącego w obwodzie bramki. Utrzymanie tyrystora w stanie przewodzenia wymaga przepływu prądu obciążenia I_F o wartości większej od dynamicznego prądu podtrzymania I_{HS} . Przepływ prądu obciążenia o takiej wartości przez tyrystor nie zanika po zaniku prądu bramki.

Właściwości zaporowe tyrystor odzyskuje po zmniejszeniu prądu głównego poniżej wartości prądu podtrzymania I_{HS} . W obwodzie prądu przemiennego wyłączenie tyrystora następuje pod wpływem zmiany kierunku przepływu prądu głównego, spowodowanej zmianą kierunku napięcia występującego między anodą i katodą. Odzyskanie przez tyrystor właściwości zaporowych wymaga czasu niezbędnego do usunięcia ładunku przestrzennego z wewnętrznego złącza PN. Czas wyłączenia w zależności od wartości prądu przewodzenia wynosi od 100 do 300 μs , jedynie dla tyrystorów szybkich nie przekracza 5 μs . Zjawisko odwrotne występuje podczas narastania napięcia anoda-katoda, gdy tyrystor znajduje się w stanie blokowania. W tym stanie napięcie anoda-katoda jest dodatnie, odkłada się na środkowym, nieprzewodzącym złączu PN i podczas narastania ładuje prądem wstecznym pojemność tego złącza. Przy zbyt dużej szybkości narastania tego napięcia może dojść do odblokowania tyrystora i przejścia do stanu przewodzenia bez udziału prądu bramki. Aby tego uniknąć, stromość narastania napięcia blokowania nie może być większa od podawanej w katalogach wartości dopuszczalnej, mieszczącej się w zakresie od 50 V/ μs do 1000 V/ μs . Podobnie jak napięcie blokowania, nie może również zbyt gwałtownie narastać prąd przewodzenia I_F tyrystora, ponieważ pastylka półprzewodnika, z której jest on wykonany, może ulec punktowemu przegrzaniu. Dopuszczalna

dla danego typu tyrystora stromość narastania prądu przewodzenia podawana jest w katalogach i wynosi od $100 \text{ A}/\mu\text{s}$ do $1000 \text{ A}/\mu\text{s}$.

Ze względu na właściwości tyrystory blokowane napięciem wstecznym mogą być stosowane nie tylko jako łączniki bezstykowe, ale również jako prostowniki sterowane.

Oprócz tyrystorów blokowanych napięciem wstecznym produkuje się również tyrystory blokowane napięciem bramki, tzw. **tyrystory GTO** (ang. **Gate Turn Off thyristor**). Udrożnienie (odblokowanie) tego tyrystora następuje za pomocą dodatniego impulsu prądu bramkowego, do zablokowania zaś wystarczy podać impuls prądu bramki o odwrotnej polaryzacji. Wyłączający impuls prądu bramki wymusza w strukturze tyrystora przepływ prądu o polaryzacji przeciwnej do prądu przewodzenia, usuwając w ten sposób ładunek przestrzenny z obszaru środkowego złącza PN. Produkowane tyrystory GTO mają napięcia wsteczne nieprzekraczające wartości 5000 V i prąd przewodzenia o wartości nie większej niż 3000 A . Częstotliwość przełączeń takich tyrystorów może osiągać wartość 10 kHz . Symbol graficzny tyrystora GTO przedstawiono na rys. 4.27a.



Rys. 4.27. Symbole graficzne: tyrystora GTO (a), układu odwrotnie równoległego (b) oraz triaka (c)

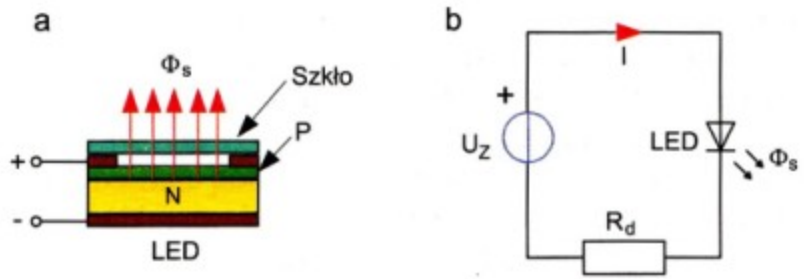
Do sterowania wartością skuteczną prądu przemiennego można wykorzystać dwa tyrystory wyłączane napięciem wstecznym, połączone odwrotnie równolegle (rys. 4.27b), sterowane odrębnie za pośrednictwem swoich bramek. Odmianą takiego układu jest **triak**, utworzony z dwóch odwrotnie równoległych tyrystorów wytworzonych w jednej strukturze półprzewodnika. Oznaczenie graficzne triaka przedstawiono na rys. 4.27c. W odróżnieniu od układu odwrotnie równoległego, triak jest wyzwalany za pomocą jednego wyprowadzenia bramkowego przy zasilaniu zarówno stałoprądowym, jak i prądu przemiennego. Ze względu na właściwości triak jest nazywany również tyrystorem symetrycznym dwukierunkowym.

4.7. Podzespoły fotoelektryczne

4.7.1 Dioda elektroluminescencyjna LED

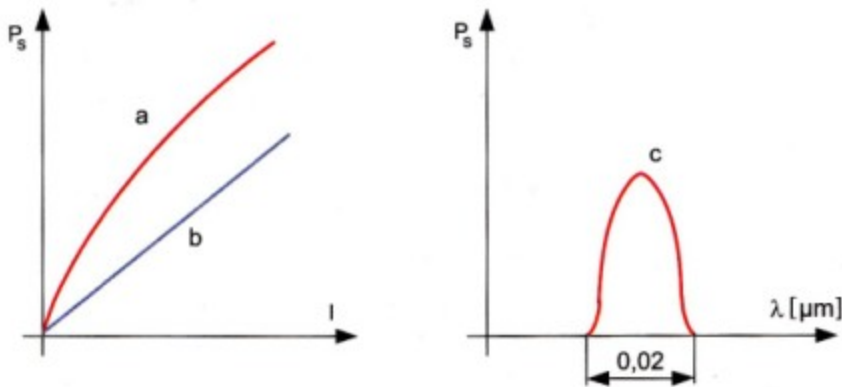
Półprzewodnikowa dioda elektroluminescencyjna LED (ang. **Light Effect Diode**), emitująca promieniowanie widzialne lub niewidzialne podczerwone pod wpływem płynącego przez nią prądu elektrycznego, jest nazywana również diodą

Rys. 4.28. Struktura wewnętrzna (a) oraz sposób zasilania diody elektroluminescencyjnej (b)



świecącą. Przepływ prądu przewodzenia przez taką diodę powoduje rekombinację nośników mniejszościowych (par dziur i elektronów) w złączu PN, któremu towarzyszy emisja fotonów. Konstrukcję oraz sposób zasilania diody LED przedstawiono na rys. 4.28.

W diodach LED występuje najczęściej zjawisko emisji spontanicznej i moc P_s emitowanego światła wzrasta proporcjonalnie do wartości prądu przewodzenia (rys. 4.29).



Rys. 4.29. Charakterystyki zależności mocy P_s emitowanej przez diody LED od prądu zasilania I : diody świecącej całą powierzchnią złącza (a) i krawędzią złącza (b) oraz charakterystyka widmowa światła emitowanego przez diodę (c)

Diody LED emitują światło monochromatyczne o widmie, którego szerokość nie przekracza $0,02 \mu\text{m}$. Kolor emitowanego światła określa materiał, z którego wykonano diodę. Złącza PN diod emitujących promieniowanie podczerwone są wykonane z arsenku galu lub fosforku indu, diod emitujących światło czerwone i zielone – z fosforku galu, a diod emitujących światło żółte – z arsenku galu.

W celu zabezpieczenia przed uszkodzeniem termicznym spowodowanym przepływem prądu o zbyt dużej wartości, diody LED muszą być dołączone do źródła napięcia za pośrednictwem rezystora R_d (rys. 4.28). Gdy znana jest wartość dopuszczalnego maksymalnego prądu I_{max} diody oraz wartość napięcia zasilającego U_z , wartość rezystancji tego opornika wyznaczamy z zależności

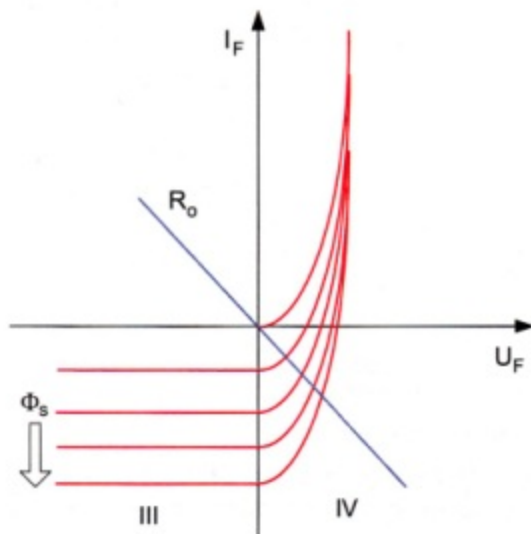
$$R_d = \frac{U_z}{I_{\text{max}}}$$

Nieuwzględnienie przy obliczeniach rezystancji opornika R_d wartości rezystancji wewnętrznej diody spolaryzowanej w kierunku przewodzenia oraz rezystancji źródła napięcia U_z stanowi dodatkowe zabezpieczenie diody.

We współczesnych samochodach diody LED zupełnie wyparły żarówki małej mocy stosowane wcześniej w lampkach sygnalizacyjnych. Diody LED wykorzystuje się w energooszczędnych źródłach światła umieszczanych w reflektorach oraz tylnych światłach zewnętrznych. W porównaniu z tradycyjnymi żarówkami z drutem grzejnym, zespoły świetlne diod LED emitują więcej światła przy mniejszym poborze energii elektrycznej. Na przykład zespół 18 diod LED o mocy 1,5 W ma wydajność świetlną większą niż żarówka samochodowa o mocy 5 W.

4.7.2. Fotodiody

Oświetlenie spolaryzowanego zaporowo złącza PN powoduje wzrost liczby nośników mniejszościowych w tym złączu. Energia świetlna powoduje generowanie par elektron-dziura, których liczba w jednostce czasu jest proporcjonalna do natężenia światła N_s . Skutkiem tego zjawiska jest kilkukrotne zwiększenie



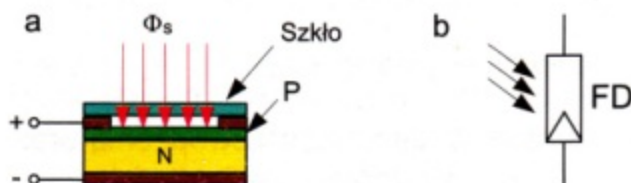
Rys. 4.30. Rodzina charakterystyk $I_F = f(U_F)$ fotodiody w zależności od natężenia oświetlenia

prądu wstecznego, który ze względu na przyczynę jest nazywany fotoprądem I_F . Zjawisko generacji dodatkowych par nośników ładunku elektrycznego w złączu PN jest nazywane efektem fotoelektrycznym wewnętrznym i ma również wpływ na przebieg charakterystyki prądowo-napięciowej fotodiody $I_F = f(U_F)$ – rys. 4.30.

Docierająca do fotodiody energia świetlna powoduje przesunięcie charakterystyki $I_F = f(U_F)$ do IV ćwiartki układu współrzędnych (rys. 4.30). W tym obszarze dioda LED jest elementem aktywnym, generującym prąd od obwodu elektrycznego. Fotodiodę w tym stanie pracy nazywa się fotoogniwem i wykorzystuje między innymi do budowy baterii słonecznych.

Do produkcji fotodiod wykorzystuje się najczęściej krzem. W porównaniu z diodami prostowniczymi złącze PN fotodiod ma znacznie większą powierzchnię, od 1 mm² do kilku cm². Budowę i symbol fotodiody przedstawiono na rys. 4.31.

Najczęściej fotodiodę wykorzystuje się do wykrywania sygnału świetlnego. W tym celu trzeba ją spolaryzować w kierunku zaporowym, tzn. jej punkt pra-



Rys. 4.31. Struktura wewnętrzna fotodiody (a) oraz jej symbol (b)

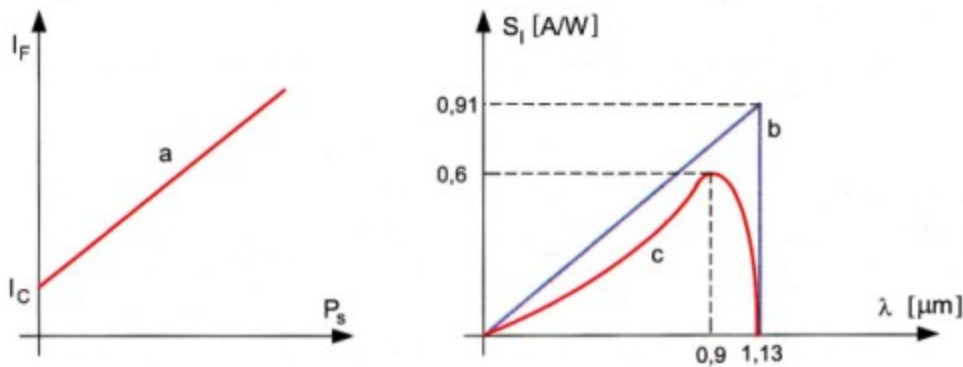
cy musi się znajdować w trzeciej ćwiartce układu współrzędnych. Przy takiej polaryzacji przez fotodiode płynie fotoprąd I_F o wartości zależnej jedynie od natężenia światła padającego na jej złącze PN i niezależnej od napięcia polaryzującego diodę. Ze względu na nieprzekraczającą kilkudziesięciu mikroamperów wartość prądu I_F , fotodioda jest połączona najczęściej z układem wzmacniającym ten prąd.

Dodatkowe pary nośników prądu elektrycznego w złączu PN fotodiody generują fotony o energii niezbędnej do ich aktywacji. W związku z tym fotodiody generują fotoprąd I_F jedynie pod wpływem światła o określonej długości fali. Właściwość ta jest charakteryzowana w katalogach za pośrednictwem tzw. czułości prądowej S_I definiowanej za pomocą zależności:

$$S_I = \frac{I_F - I_C}{P_s} = \frac{I_s}{P_s} \left[\frac{\text{A}}{\text{W}} \right] \quad (4.6)$$

W zależności tej I_C jest wstecznym prądem ciemnym o wartości nieprzekraczającej kilku miliamperów, płynącym przez nieoświetlone złącze PN diody pod wpływem temperatury otoczenia. I_F jest fotoprądem diody, którego przepływ jest uwarunkowany oświetleniem złącza PN promieniowaniem o mocy P_s . Jednostką czułości fotodiody jest wartość prądu płynącego przez oświetlone złącze, przypadająca na jednostkę mocy promieniowania świetlnego, tzn. amper na wat.

Prąd I_F wzrasta proporcjonalnie do mocy promieniowania P_s , zgodnie z charakterystyką prądowo-oświetleniową przedstawioną na rys. 4.32a.



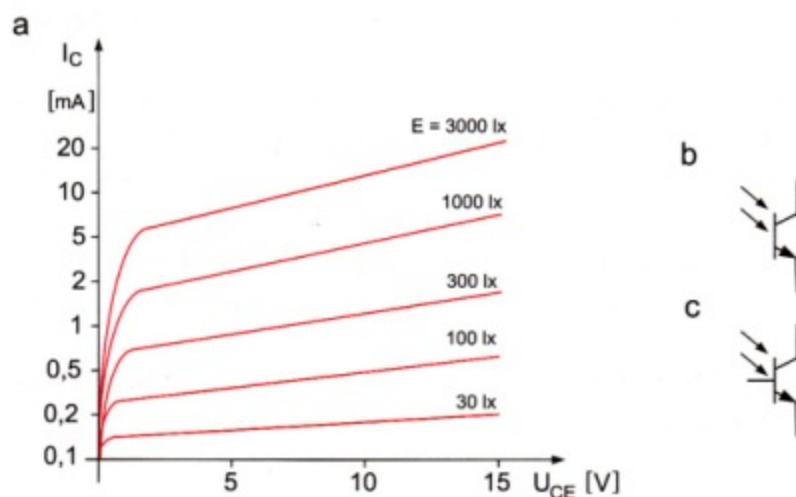
Rys. 4.32. Charakterystyka prądowo-oświetleniowa fotodiody krzemowej (a) oraz jej czułość w funkcji długości fali światła padającego: idealna (b) i rzeczywista (c)

Zależność czułości fotodiody od długości fali świetlnej docierającej do złącza PN, tzn. charakterystykę $S_I = f(\lambda)$, przedstawiono na rys. 4.32b. Z przebiegu tej charakterystyki wynika, że parametr S_I fotodiody przyjmuje wartość maksymalną zbliżoną do 0,6 A/W dla $\lambda = 0,9 \mu\text{m}$.

Fotodiody wykorzystuje się w opisywanych dalej transoptorach oraz jako odbiorniki w transmisji sygnałów za pomocą światłowodów, stanowiących warstwę fizyczną magistrali komunikacyjnej MOST, która łączy podzespoły samochodowych zestawów multimedialnych.

4.7.3. Fototranzystor

Z użytkowego punktu widzenia działanie fototranzystorów bipolarnych jest podobne do działania fotodiod. Światłne promieniowanie widzialne lub niewidzialne, docierając do spolaryzowanego zaporowo złącza PN kolektor-baza, generuje dodatkowe pary elektron-dziura i wywołuje przepływ fotoelektrycznego prądu kolektora o wartości proporcjonalnej do natężenia promieniowania oświetlającego fototranzystor. Wpływ promieniowania świetlnego na prąd kolektorowy ilustrują charakterystyki wyjściowe $I_{CF} = f(U_{CE})$ fototranzystora – rys. 4.33. Na rysunku tym przedstawiono również symbole graficzne fototranzystora.



Rys. 4.33. Charakterystyka $I_{CF} = f(U_{CE})$ fototranzystora (a) oraz symbol fototranzystora bez wyprowadzonej bazy (b) i z wyprowadzoną bazą (c)

Działanie promieniowania świetlnego na złącze kolektor-emiter fototranzystora można porównać z przepływem dodatkowego prądu w obwodzie bazy. Ze względu na wzmacnianie prądowe fototranzystory mają większą czułość niż fotodiody.

Fototranzystory są sterowane promieniowaniem świetlnym i dlatego do ich działania nie jest wymagany przepływ prądu w obwodzie bazy. W takiej konfiguracji pracują fototranzystory sterowane impulsami świetlnymi. W odniesieniu do fototranzystorów przeznaczonych do wykrywania płynnych zmian natężenia światła jest niezbędne ustalenie i stabilizacja punktu pracy obwodu bazy i dlatego takie fototranzystory są wyposażone w wyprowadzenie bazy.

Z konstrukcji fototranzystorów wynika duża pojemność złącza kolektor-emiter, która wydłuża czas przełączania ze stanu przewodzenia do stanu zaporowego, dlatego mogą one być stosowane w zakresie częstotliwości nie większych niż 250 kHz.

Typowym przykładem praktycznego zastosowania fototranzystorów są opisane dalej transoptory.

4.7.4. Transoptor

Transoptor jest elementem optoelektronicznym składającym się z nadajnika i odbiornika promieniowania umieszczonych w jednej, nieprzepuszczającej światła obudowie. W literaturze nadajnik i odbiornik promieniowania nazywa się odpowiednio fotoemiterem oraz fotodetektorem. Symbol graficzny transoptora oraz sposób połączenia tego podzespołu do elektrycznych układów zewnętrznych przedstawiono na rys. 4.34.

W zależności od przeznaczenia transoptora sprzężenie optyczne między fotoemiterem oraz fotodetektorem umożliwia powietrze lub światłowód wykonany z tworzywa sztucznego, szkła albo przezroczystej żywicy epoksydowej.

W transoptorach rolę nadajnika odgrywa dioda elektroluminescencyjna, emitująca najczęściej niewidzialne promieniowanie podczerwone, a fotodetektorem jest fotodioda lub fototranzystor.

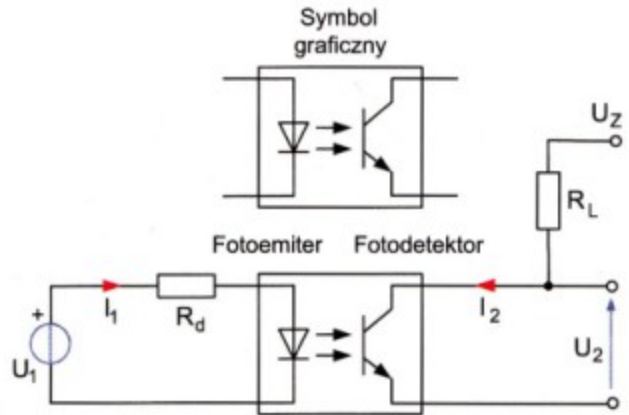
Działanie transoptora jest następujące (rys. 4.34). Przepływ prądu I_1 jest przyczyną promieniowania diody LED docierającego do fotodetektora, którym w naszym przypadku jest fototranzystor. Skutkiem tego promieniowania jest przepływ prądu I_2 fototranzystora, o wartości proporcjonalnej do prądu I_1 . Parametrem transoptorów określającym zależność prądu wyjściowego I_2 od prądu wejściowego I_1 jest **współczynnik przenoszenia prądu stałego CTR** (ang. *Current Transfer Ratio*), zdefiniowany zależnością:

$$\text{CTR} = \frac{I_2}{I_1} \quad (4.7)$$

Wartość współczynnika CTR transoptorów z fotodiodą nie przekracza 0,002, a transoptorów z fotodetektorem tranzystorowym wynosi od 10 do 500.

Małą wartość współczynnika CTR transoptory z fotodiodą rekompensują bardzo krótkim czasem przełączania – rzędu ns. Dzięki temu mogą przenosić sygnały o częstotliwości do 10 MHz. W odróżnieniu od transoptorów z fotodiodami, transoptory z fototranzystorami charakteryzują się długim czasem przełączania i przez to częstotliwość ograniczająca od góry ich pasmo przenoszenia ma wartość około 500 kHz.

W transoptorach obwód wejściowy oraz wyjściowy są rozdzielone galwanicznie przerwą izolacyjną o rezystancji rzędu $10^{11} \Omega$. Napięcie probiercze, przy którym rezystancja przerwy galwanicznej separuje oba układy, ma wartość kilku kilowoltów.



Rys. 4.34. Symbol graficzny oraz sposób zasilania transoptora

Transoptory są stosowane do separacji galwanicznej obwodów różniących się napięciem pracy. Przenoszą sygnały zarówno stałe, jak i przemienne. Wymienione funkcje transoptory realizują również w technice samochodowej, gdzie dodatkowo stanowią element przetworników wielkości elektrycznych, takich jak przemieszczenie liniowe lub kątowe, na sygnał elektryczny.

4.8. Pytania i zadania

1. Jakimi właściwościami charakteryzują się materiały półprzewodnikowe?
2. Czym różnią się półprzewodniki typu N od półprzewodników typu P?
3. Jaka różnica występuje między półprzewodnikiem samoistnym a niesamoistnym?
4. Co to jest złącze półprzewodnikowe i jakie zjawiska zachodzą w tym złączu?
5. W jakich stanach może pracować dioda półprzewodnikowa?
6. Jaki przebieg ma charakterystyka prądowo-napięciowa diody prostowniczej?
7. Jak należy spolaryzować złącze PN diody Zenera, aby stabilizowała ona napięcie?
8. Co to jest współczynnik stabilizacji diody Zenera?
9. Wyjaśnij zasadę działania diody pojemnościowej.
10. Jak działa tranzystor bipolarny?
11. Jaka jest różnica między przebiegiem charakterystyki wejściowej tranzystora bipolarnego i diody półprzewodnikowej?
12. Wykorzystując charakterystykę wyjściową, opisz stany pracy tranzystora bipolarnego.
13. Wyjaśnij wzmacniające działanie tranzystora bipolarnego.
14. Podaj zależność definiującą współczynnik wzmocnienia prądowego tranzystora bipolarnego.
15. Wyjaśnij zasadę działania tranzystora polowego złączowego JFET.
16. Czy tranzystor polowy złączowy jest sterowany napięciem, czy prądem bramki?
17. Do czego można wykorzystać charakterystykę przejściową tranzystora JFET?
18. Co to są charakterystyki wyjściowe tranzystora JFET?
19. Wyjaśnij działanie tranzystorów polowych z izolowaną bramką.
20. Jaka wartość przyjmuje prąd kanału tranzystora D-MOS, gdy $U_{GS} = 0$?
21. Jaka wartość przyjmuje prąd kanału tranzystora E-MOS, gdy $U_{GS} = 0$?
22. Wyjaśnij działanie tyrystora, wykorzystując jego charakterystykę prądowo-napięciową.
23. Jak można przerwać przepływ prądu przez tyrystor?
24. Wyjaśnij różnice w działaniu fotodiody oraz diody LED.

25. Zaprojektuj układ sygnalizujący obecność napięcia stałego o wartości 12 V, składający się z szeregowo połączonych diody prostowniczej D , opornika R_s oraz diody LED. Dioda LED pracuje prawidłowo, gdy płynie przez nią prąd o wartości nieprzekraczającej 30 mA. Przy takim prądzie na diodzie występuje spadek napięcia o wartości 2 V.

Wskazówka: sposób rozwiązania zadania opisano w podrozdziale 4.7.1.

26. Opisz budowę i działanie fototranzystora.
27. Czy fototranzystory mają wyprowadzenie bazy?
28. Jak działa transoptor?
29. Jaką zależnością jest zdefiniowany współczynnik przenoszenia prądu stałego przez transoptor?
30. Dlaczego transoptor separuje galwanicznie dwa obwody?