

tronów, są podstawą do przypuszczeń, że trudności te zostaną przezwyciężone w znacznym stopniu zmniejszone. W niniejszej książce jest zapoznanie czytelnika z zasadą działania półprzewodnikowych, ich konstrukcją i własnościami. Niektórych fragment poświęcono także laserom złączowym sprężonym. Te ostatnie mogą być w przyszłości przydatne w dziedzinie matematycznych opartych na technice świetlnej. Związkiem, który jak dotąd znalazł najszersze zastosowanie w logii laserów złączowych, jest arsenek galu. Fakt ten należy do zarówno własnościom fizycznym arsenku galu, szczególnie nim do wymuszenia akcji laserowej, jak również stosunkowo panowanej technologii syntezy oraz monokryształizacji tego półprzewodnika. Dlatego też przeważająca część materiału zawartego w książce będzie dotyczyła laserów wykonanych z arsenku galu.

1. PODSTAWOWE POJĘCIA Z ZAKRESU ELEKTRONIKI KWANTOWEJ

1.1. Pojęcie spójności promieniowania

Światło rozpatrywane jako promieniowanie elektromagnetyczne charakteryzuje się natężeniem, częstotliwością oraz fazą i polaryzacją składających się na nie fal. O ile dwie pierwsze cechy sprostęga się stosunkowo łatwo np. za pomocą oka, jeśli widmo leży w pasmie widzialnym, o tyle różnice w fazie możemy stwierdzić dopiero drogą specjalnych doświadczeń. Zwykle polegają one na badaniu interferencji między falami wysyłanymi z różnych punktów źródła lub z tego samego punktu, lecz w różnych chwilach czasu. Interferencja jest wynikiem wektorowego sumowania się ciągów fal i jest dowodem istnienia korelacji między fazą drgań w poszczególnych ciągach. Promieniowanie wykazujące taką korelację fazową nazywa się *spójnym*¹⁾.

Doświadczenia przeprowadzane z klasycznymi źródłami światła wykazują, że nawet w przypadku gdy emitowane promieniowanie ma bardzo wąskie widmo, a więc jest prawie monochromatyczne, spójność występuje jedynie między ciągami fal rozchodzącymi się dostatecznie blisko siebie w sensie przestrzeni i czasu. Dla jasności mówię się zatem o spójności przestrzennej i spójności czasowej.

Przez pojęcie spójności przestrzennej rozumiemy stopień korelacji między fazami promieniowania monochromatycznego emitowanego przez dwa różne punkty źródła o skończonych wymiarach. Miara tej korelacji jest zdolność do interferencji.

Spójność czasową można ocenić badając interferencję między ciągami fal wychodzącymi z jednego punktu źródła, lecz przebiegającymi drogami o różnej długości. Interferencja przestaje występować po przekroczeniu pewnej różnicy długości dróg takiej, że w przedziale czasu Δt , potrzebnym do jej pokonania, nastąpiła przypadkowa zmiana fazy. Ten kry-

¹⁾ Pojęcie korelacji zostało tu użyte w powszechnie przyjętym dla tego słowa znaczeniu odpowiedniości. Bardziej precyzyjne omówienie zagadnienia spójności promieniowania wykracza poza ramy niniejszej książki. Można je znaleźć np. w pracy Borna i Wolfa [1.2].

przedziału czasu Δt określa właśnie spójność czasową, tzn. dla $\Delta t \ll \Delta t$ promieniowanie można uważać za prawie spójne. Brak z kolei fazy w czasie oznacza jednocześnie, że promieniowanie nie jest chromatyczne. Szerokość jego widma $\Delta\nu$ jest odwrotnie proporcjonalna do przedziału czasu Δt określającego spójność czasową. Użyte wyrażenie „prawie spójne” ma na celu podkreślenie faktu, że wyrodzie nie występują źródła emitujące promieniowanie idealnie podobnie jak nie ma źródeł światła ściśle monochromatycznego.

1.2. Zjawiska absorpcji emisji i wzmacniania

odnie z podstawowym założeniem fizyki kwantowej atomy oraz cząstki mogą pozostawać jedynie w określonych stanach stacjonarnych, z których każdy odpowiada określonej wartości energii nazywanej *poziomem układu atomowego*. Może się zdarzyć, że jednemu poziomowi odpowiada więcej stanów niż jeden, występuje wówczas tzw. *degeneracja poziomu*, a liczba stanów o tej samej energii nazywa się *liczbą poziomów*. Obsadzenie poszczególnych poziomów w stanie równowagi jest określone przez prawo Boltzmann'a, które dla układu zamkniętego w postaci

$$\frac{N_n}{g_n} = \frac{N_m}{g_m} \exp \left(-\frac{E_n - E_m}{kT} \right) \quad (1.1)$$

oznacza:

N_n — liczba atomów w stanie n i m ;

E_n — energie tych stanów;

g_n — krotność poziomów n i m .

Wzrost liczby atomów w stanie n może nastąpić w wyniku absorpcji energii E_n lub jej emisji, np. w postaci promieniowania. Wzrost liczby atomów w stanie m może nastąpić w wyniku emisji energii E_m .

Wzrost liczby atomów w stanie n może nastąpić również w wyniku absorpcji energii lub jej emisji, np. w postaci promieniowania. Wzrost liczby atomów w stanie m może nastąpić w wyniku emisji energii E_m .

$$h\nu = E_n - E_m \quad (1.2)$$

oznacza:

h — stała Plancka;

ν — częstotliwość promieniowania.

Wzrost liczby atomów w stanie n może nastąpić w wyniku absorpcji energii E_n lub jej emisji, np. w postaci promieniowania. Wzrost liczby atomów w stanie m może nastąpić w wyniku emisji energii E_m .

$$P_{nm} = u_n B_{nm} \quad (1.3)$$

Uzyskany w wyniku absorpcji rozkład energii atomów przestał być zgodny z prawem Boltzmann'a; uważa się wówczas, że układ atomów znajduje się w stanie wzburzonym. Układ taki może przejść bez zewnętrznej przyczyny do stanu podstawowego, oddając swą energię w postaci emisji promieniowania, zwanej *emisją spontaniczną*. Prawdopodobieństwo przejścia atomu do stanu niższego w wyniku emisji spontanicznej zostanie oznaczone przez A_{nm} . Jeśli przez N_n wyrazi się liczbę atomów na n -tym poziomie, to całkowita liczba przejść w czasie jednej sekundy z poziomu n na poziom m będzie równa $N_n A_{nm}$, a emitowana moc promieniowania wyniesie $N_n(E_n - E_m)A_{nm}$.

Przejście atomu do poziomu niższego może być również zainicjowane np. przez promieniowanie elektromagnetyczne o odpowiedniej częstotliwości określonej przez warunek (1.2). Występuje wówczas zjawisko *emisji wymuszonej*, zwanej też *niekiedy stymulowaną*. Jest przy tym rzeczą charakterystyczną i zasadniczej wagi, że zarówno częstotliwość jak i faza promieniowania pochodzącego od emisji wymuszonej będą w tym przypadku identyczne z częstotliwością i fazą promieniowania wymuszającego.

Jeżeli prawdopodobieństwo wystąpienia emisji wymuszonej zostanie oznaczone przez B_{nm} , a gęstość promieniowania wymuszającego przez u_n , to prawdopodobieństwo przejścia układu ze stanu n do m w wyniku emisji wymuszonej będzie iloczynem tych dwóch wielkości, tzn. $u_n B_{nm}$.

Miedzy wymienionymi wyżej wielkościami zachodzą zależności zwane związkami Einsteina, które mają postać:

$$\begin{aligned} g_n B_{nm} &= g_m B_{mn} \\ A_{nm} &= \frac{8\pi h \nu_{nm}^3}{c^3} B_{nm} \end{aligned} \quad (1.4)$$

przy czym:

c — prędkość światła;

ν_{nm} — współczynnik załamania ośrodka.

Teraz zostanie rozpatrzona sytuacja, w której na układ atomów mających stany podstawowy i pobudzony działa wiązka promieniowania o częstotliwości ν_{nm} , odpowiadającej różnicy energii tych stanów. Liczba przejść w dół, tzn. emisyjnych, wykonanych w czasie jednej sekundy, będzie równa sumie

$$A_{nm} N_n + u_n B_{nm} N_n \quad (1.5)$$

a liczba przejść w górę (absorpcyjnych) wyniesie odpowiednio

$$u_n B_{mn} N_m \quad (1.6)$$

Jeśli $N_n < N_m$ dla $n > m$, co występuje dla stanu równowagi, to rezultatem końcowym opisanego procesu będzie zmniejszenie intensywności wiązki wymuszającej, na korzyść emisji spontanicznej, która niekoniecznie musi zachodzić bezpośrednio między poziomami n a m .

Należenie promieniowania przechodzącego przez rozpatrywany ośrodek jest funkcją częstotliwości promieniowania oraz przebytej przez niego odległości i może być określone za pomocą wzoru

$$I = I_0 \exp(-kx) \quad (1.7)$$

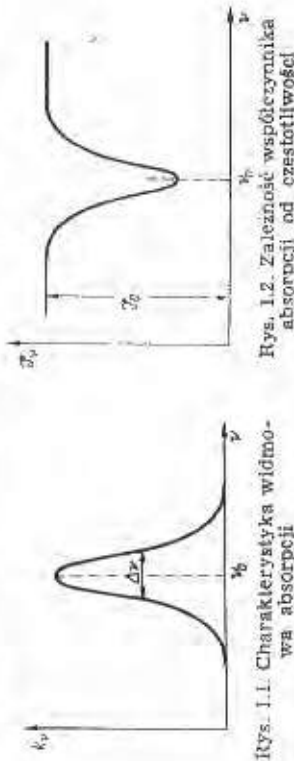
którym:

I_0 — natężenie promieniowania na powierzchni ośrodka;

k — współczynnik absorpcji będący funkcją częstotliwości;

x — odległość liczona od powierzchni.

Wykreślając tę funkcję dla ustalonej wartości x otrzymuje się w praktyce krzywą jak na rys. 1.1. Wyznaczając na jej podstawie współczynnik absorpcji jako funkcję częstotliwości, dostaje się krzywą pokazaną na



Rys. 1.1. Charakterystyka widmowa absorpcji

Rys. 1.2. Zależność współczynnika absorpcji od częstotliwości

s. 1.2. Można wykazać, że powierzchnia ograniczona przez tę krzywą może być w sposób jednoznaczny ze współczynnikami Einsteina i zatem ta ma postać [1.1]

$$\int k_p d\nu = \frac{c^2 A_{nm}}{8\pi\nu_0^2} \frac{g_n}{g_m} \left(N_m - \frac{g_m}{g_n} N_n \right) \quad (1.8)$$

Ponieważ średni czas życia τ atomu w stanie pobudzonym jest odrotnie proporcjonalny do prawdopodobieństwa przejścia, współczynnik A_{nm} w tym wzorze można zastąpić przez $\frac{1}{\tau}$ i w tej postaci będzie wykorzystany w jednym z następnych rozdziałów.

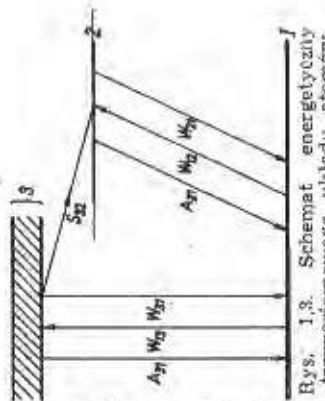
Tenże rozdział rozpatrzy układ, w którym liczba atomów wzbudzonych jest większa od liczby atomów w stanie podstawowym, tzn. chodzi o nierówność $N_n > N_m$ dla $n > m$. Sytuacja taka nazywa się wersją obsadzoną. Rozumując w sposób analogiczny do wyżej podanego, dochodzi się do wniosku, że w układzie z inwersją obsadzeń występuje absorpcja ujemna, a więc wzmocnienie wiązki promieniowania wymuszającego. Wzmocnienie to jest konsekwencją przewagi promieniowania wymuszonego nad promieniowaniem absorbowanym. W materiale, w którym jest spełniony warunek absorpcji ujemnej w pewnym zakresie częstotliwości padającego światła, natężenie promieniowania

będzie wzrastać zgodnie z funkcją (1.7) po wstawieniu do jej wykładnika współczynnika wzmocnienia

$$a = -k, \quad (1.9)$$

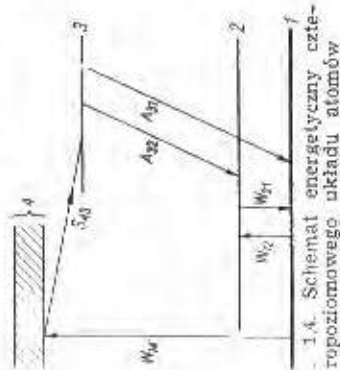
Na możliwości uzyskania wzmocnienia promieniowania w wyniku emisji wymuszonej została oparta koncepcja przyrządu zwanego laserem od angielskiego akronimu, czyli połączenia pierwszych liter nazwy angielskiej Light Amplification by Stimulated Emission of Radiation. Na marginesie warto dodać, że laser jest generatorem promieniowania, a nie jego wzmacniaczem, jak mogłoby to sugerować nazwa przyrządu.

W celu zapoznania się z mechanizmem działania lasera zostanie rozpatrzony układ atomów, których poziomy energetyczne są usytuowane w sposób przedstawiony na rys. 1.3, przy czym czas życia atomu na poziomie 2 jest znacznie dłuższy niż na poziomie 3. Pod wpływem zewnętrznej przyczyny atomy mogą być pobudzone np. w taki sposób, że ich elektrony przejdą z poziomu 1 na poziom 3. Przejście to oznaczono przez W_{13} . Stan taki jest oczywiście nietrwały i po krótkim czasie określonym przez czas życia część elektronów „spadnie” na poziom 1, czemu towarzyszy emisja spontaniczna A_{31} oraz wymuszona W_{31} . Pozostała część elektronów może jednak w wyniku emisji spontanicznej niepromienistej S_{32} przejść na poziom 2. Czas życia atomu na tym poziomie jest, jak powiedziano wyżej, dłuższy niż na poziomie 3, wystąpi zatem pewnego rodzaju akumulacja elektronów na tym poziomie, chociaż będą one również spadały na poziom 1 w wyniku emisji wymuszonej W_{21} i spontanicznej A_{21} . Emisji tej będzie towarzyszyć promieniowanie. Niezależnie od wymienionych przejść, przez cały czas również będzie zachodzić absorpcja wymuszona W_{12} z poziomu 1 na poziom 2. Jeżeli są spełnione odpowiednie warunki na szybkości przejść między poszczególnymi poziomami, a pobudzanie odbywa się dostatecznie intensywnie, to w opisanej sytuacji może wystąpić inwersja obsadzenia poziomu 2 względem poziomu 1, w wyniku czego wystąpi wzmocnienie promieniowania wymuszonego — jak stwierdzono to wyżej. Promieniowanie to rozchodzi się w ośrodku czynnym we wszystkich kierunkach i na ogół opuszcza ten ośrodek, nie mając wpływu na dalszy przebieg zjawiska. Jeżeli jednak stworzy się sprzężenie zwrotne, tzn. spowoduje powrót części energii do ośrodka czynnego, np. przez odbicie od odpowiednio ustawionych zwierciadeł, to promieniowanie o określonej energii i kierunku



Rys. 1.3. Schemat energetyczny trypozomowego układu atomów

zmieć ponownie udział w procesie wymuszania emisji. W warunkach, jeśli wzmocnienie dostatecznie przewyższy straty, może dojść do budzenia się drgań w układzie. Częstotliwość tych drgań będzie narzucona przez odległość między poziomami 2 a 1 oraz własności wnęki rezonansowej, utworzonej przez osrodek czynny i zwierciadła. Jest rzeczą intuicyjnie wyczuwalną, że promieniowanie to będzie prawie monochromatyczne i jego fazy uporządkowane w takim stopniu, że można zacząć je za spójne. Proces generacji drgań w laserze określa się niekiedy terminem *akcja laserowa*, a moment ich wzbudzenia *przekroczeniem progu pobudzenia lasera*. Konsekwentnie, stan lasera, w którym występuje akcja laserowa, jest często nazywany *pobudzeniem lasera*.



1.4. Schemat energetyczny cztero-poziomowego układu atomów

Opisany wyżej laser nosi nazwę *trójpłaszczyznowego* z uwagi na to, że w akcji laserowej biorą udział 3 poziomy. Do grupy laserów tego typu należą np. lasery wykonane z rubinu domieszkowanego chromem. Zasadniczą wadą lasera trójpłaszczyznowego jest konieczność wywołania inwersji obsadzeń w stosunku do poziomu podstawowego, z natury rzeczy obserwowanego. Znacznie łatwiej byłoby ją osiągnąć w układzie przedstawionym na rys. 1.4, w którym pożądana emisja wymuszona zachodzi między poziomami 3 i 2, podczas gdy źródłem elektronów jest atom podstawowy 1. Układ taki został w praktyce zrealizowany między innymi w kryształach fluorku wapnia domieszkowanym jonami uranu oraz mieszaninie hel-neon. Lasery pracujące na tej zasadzie noszą nazwę *czteropłaszczyznowych*. Mają one znaczną przewagę nad laserami trójpłaszczyznowymi pod warunkiem, że odległość między poziomami 1 i 2 jest dostatecznie duża, a czas życia atomu na poziomie 2 znacznie krótszy od czasu życia atomu, na poziomie 3.

2. ZASADA DZIAŁANIA LASERÓW ZŁĄCZOWYCH

2.1. Inwersja obsadzeń w półprzewodnikach

Opisane w poprzednim rozdziale przejścia optyczne, w wyniku których zachodzi emisja światła spójnego, występowały między poziomami energetycznymi o dokładnie określonej energii. W sytuacji tej miarą inwersji obsadzeń poziomów był stosunek całkowitej liczby atomów w stanie wyższym do liczby atomów w stanie niższym. W przypadku półprzewodników zagadnienie to znacznie się komplikuje, gdyż stany osiągalne dla elektronów są tutaj zgrupowane w pasmach rozciągających się w stosunkowo dużym zakresie energii. Na rys. 2.1 przedstawiono obraz takich pasm dla półprzewodnika domieszkowanego niewielką ilością akceptorów i donorów. Linia przerywaną zaznaczono rozkład gęstości stanów $N(E)$ w poszczególnych pasmach. Należy zwrócić uwagę, że obraz ten jest słuszny jedynie dla małych koncentracji domieszek. Ze wzrostem koncentracji pasma domieszkowe przesuwają się, a następnie całkowicie się zlewają z pasmami głównymi, w wyniku czego zmienia się całkowity rozkład gęstości stanów oraz szerokość pasma zabronionego.

Prawdopodobieństwo obsadzenia przez elektron poziomu o energii E jest opisane przez statystykę Fermiego-Diraca i dla stanu równowagi termodynamicznej wyraża się funkcją

$$f_n(E) = \frac{1}{1 + \exp\left(\frac{E - F}{kT}\right)} \quad (2.1)$$

przy czym:

k — stała Boltzmanna;

T — temperatura bezwzględna;

F — parametr o wymiarze energii, noszący nazwę poziomu Fermiego.

Przebieg funkcji $f_n(E)$ dla temperatury 0°K i wyżej przedstawiono na rys. 2.2. Można również mówić o prawdopodobieństwie nieobsadzenia poziomu przez elektron, które dla pasma walencyjnego będzie oznaczało prawdopodobieństwo obsadzenia stanu przez dziurę. Wyrazi się ono wzorem

$$f_p(E) = 1 - f_n(E) = \frac{1}{1 + \exp\left(\frac{F - E}{kT}\right)} \quad (2.2)$$

Ilość gęstości stanów $N(E)$ i prawdopodobieństwa ich obsadzenia $f_n(E)$ wyznacza liczbę elektronów dn zajmujących poziomy w przedziale energii dE

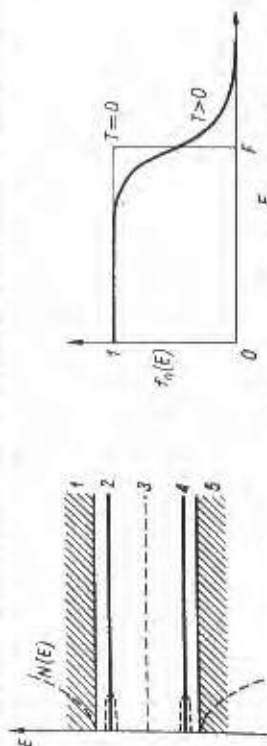
$$dn = f_n(E) N(E) dE \quad (2.3)$$

Całkowita liczba elektronów w pasmie będzie całką z tego wyrażenia zciągniętą na wszystkie poziomy pasma

$$n = \int f_n(E) N(E) dE \quad (2.4)$$

Równanie to umożliwia znalezienie poziomu Fermiego, jeśli jest znana zba elektronów n .

Termodynamiczna równowaga półprzewodnika może być zaburzona zez przyczynę zewnętrzną, np. w wyniku wstrząśnięcia do niego noś-



Rys. 2.1. Poziomy energetyczny w półprzewodniku domieszkowanym donorami i akceptorami: 1 — pasmo przewodnictwa; 2 — poziomy donatorów; 3 — poziom Fermiego; 4 — poziomy akceptorowy; 5 — pasmo walencyjne

ów mniejszościowych. Sytuacja taka może wystąpić w bezpośrednim sąsiedztwie złącza $p-n$, do którego przyłożono napięcie w kierunku przewodzenia. Jeśli między tymi wstrzykniętymi nośnikami ustali się równowaga cieplna w czasie krótszym od czasu życia elektronu, to prawdopodobieństwo, że elektron będzie zajmował dany stan w pasmie przewodnictwa, może być określone przez wyrażenie podobne do (2.1) z tą różnicą, że parametr F zostanie zastąpiony przez wielkość F_c odpowiadającą położeniu poziomu zwanego poziomem quasi-Fermiego

$$f_{nc}(E) = \frac{1}{1 + \exp\left(\frac{E - F_c}{kT}\right)} \quad (2.5)$$

Lokalizacja poziomu F_c na skali energii jest ściśle związana z koncentracją elektronów wstrzykniętych do pasma przewodnictwa i może być określona na podstawie zależności

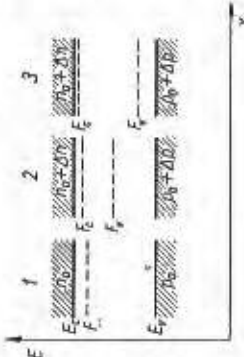
$$F_c = F + kT \ln\left(1 + \frac{\Delta n}{n_0}\right) \quad (2.6)$$

w której:

Δn — liczba wstrzykniętych elektronów do jednostkowej objętości;
 n_0 — koncentracja elektronów w pasmie przewodnictwa w stanie równowagi.

Podobne zależności występują dla elektronów w pasmie walencyjnym, ich poziom quasi-Fermiego oznaczony zostanie przez F_v .

W celu wyjaśnienia na rys. 2.3 przedstawiono schematycznie usytuowanie poziomów F_c i F_v w modelu pasmowym półprzewodnika typu n .

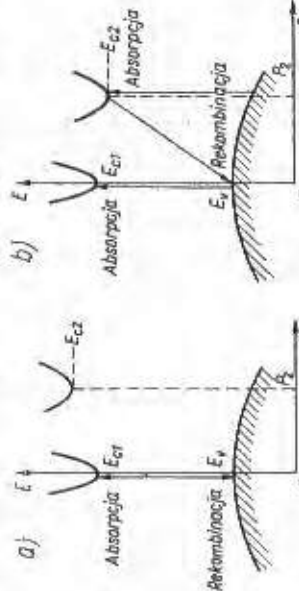


Rys. 2.3. Usytuowanie poziomów quasi-Fermiego dla różnych koncentracji wstrzykniętych nośników: 1 — równowaga; 2 — koncentracja niewielka; 3 — koncentracja duża

do którego wstrzyknięto nośniki nadmiarowe o koncentracji odpowiednio Δn i Δp . Jest rzecz intuicyjnie wyczuwalna, że usytuowanie to świadczy o stopniu, w jakim zostały obsadzone przez elektrony poziomy pasma walencyjnego i przewodnictwa, a więc może służyć jako miara inwersji obsadzenia poziomów w tych pasmach.

2.2. Rekombinacja promienista w półprzewodnikach

Proces powrotu układu do stanu stacjonarnego po zaniku czynnika zaburzającego lub czynnika utrzymującego układ w poprzednim stanie nazywa się ogólnie relaksacją. Szczególnym przypadkiem relaksacji jest



Rys. 2.4. Pasma energetyczne w półprzewodniku przedstawione jako funkcja pędu: a) półprzewodnik bezpośredni; b) półprzewodnik pośredni

rekombinacja swobodnych nośników ładunku, elektronów i dziur w półprzewodniku. Sposób oddawania energii, jaką nośnik nabył przy pobudzeniu, decyduje o rodzaju rekombinacji i z tego punktu widzenia można to można podzielić na dwie zasadnicze grupy: rekombinację promienistą i rekombinację bezpromienistą. Do ostatniej zalicza się rekombinację fononową, zderzeniową, ekscytionową, plazmową i inne.

Rodzaj rekombinacji zachodzącej w półprzewodniku ściśle wiąże się budową jego modelu pasmowego. Na rys. 2.4 przedstawiono fragmenty takiego modelu we współrzędnych energia-pęd, dla dwóch różnych półprzewodników. Zostały one uzyskane w wyniku analizy rozwiązań równania Schrödingera dla elektronów w polu potencjału sieci krystalicznej. Linie górne określają poziomy pasma przewodnictwa, dolne — poziomy pasma walencyjnego. Obszar zawarty między nimi odpowiada przerwie energetycznej. W otoczeniu punktów ekstremalnych, pasma te są opisywane za pomocą funkcji:

$$\left. \begin{aligned} E - E_{c1} &= \frac{p^2}{2m_c} \\ E - E_{c2} &= \frac{(P - P_2)^2}{2m_{c2}} \\ E - E_v &= \frac{p^2}{2m_v} \end{aligned} \right\} \quad (2.7)$$

przy czym m_c i m_v — odpowiednio masy efektywne elektronów w pasmie przewodnictwa i dziur w pasmie walencyjnym.

Istotna różnica między modelem z rys. 2.4a a modelem z rys. 2.4b polega na wartości energii odpowiadającej minimum pasma przewodnictwa dla $P = P_2$. W modelu z rys. 2.4a $E_{c1} < E_{c2}$, podczas gdy dla modelu z rys. 2.4b $E_{c1} > E_{c2}$. Oznacza to, że elektrony, które zgodnie ze statystyką Fermiego-Diraca dążą do obsadzania poziomów położonych możliwie nisko, będą wypełniały minimum E_{c1} w przypadku pierwszym i minimum E_{c2} w przypadku drugim. Przebiegi między pasmami przewodnictwa a pasmem walencyjnym, możliwe w obydwu przypadkach, zaznaczono na rys. 2.4 strzałkami. Przebiegi te podlegają regule wyboru mechaniki kwantowej. Jedną z nich — zasada Pauliego — wymaga, aby poziom końcowy, na który mógłby spaść elektron, był niezajęty, co eliminuje w procesie rekombinacji poziomy pasma walencyjnego oddalone od punktu $P=0$. Druga reguła dotyczy przejść promienistych i wymaga spełnienia warunku

$$P_c - P_v = \Delta P = 0 \quad (2.8)$$

czyli czym P_c i P_v — pęd elektronu odpowiednio w pasmie przewodnictwa i pasmie walencyjnym.

Oznacza to, że ewentualne przejścia promieniste w półprzewodniku rys. 2.4b powinny być połączone z jednoczesną zmianą pędu wywołaną

emisją lub absorpcją fononu. Przejścia takie nazywają się *pośrednimi*, a półprzewodnik, w którym występują — *półprzewodnikiem pośrednim*. Do półprzewodników takich należy np. german i krzem. W odróżnieniu od przejść pośrednich, przejścia z rys. 2.4a są nazywane *bezpośrednimi*, a towarzyszące im zjawisko emisji fotonów nosi nazwę *elektroluminescencji*. Przejścia te zachodzą w półprzewodnikach zwanych *bepośrednimi*, do których należy większość związków międzymetalicznych typu $Au-BV$ jak arsenek galu, arsenek indu, antymonek indu i inne. Można łatwo wywnioskować, że przejścia pośrednie będą znacznie mniej prawdopodobne od bezpośrednich, a zatem czas życia nośnika, który jest wielkością proporcjonalną do prawdopodobieństwa przejścia, powinien być zawsze większy w półprzewodniku pośrednim niż bezpośrednim. Wniosek ten znajduje potwierdzenie w praktyce, jednak obserwowane różnice wartości czasów życia są znacznie mniejsze, niż wynikałoby to z obliczeń teoretycznych. W celu wyjaśnienia tej rozbieżności Shockley i Read [2.1] przyjęli, że w półprzewodnikach istnieją pewne defekty sieci lub domieszki, nazywane *centrami rekombinacji*, którym odpowiadają poziomy zlokalizowane w pasmie zabronionym. Centra te katalizują proces rekombinacji dzięki temu, że może ona zachodzić w dwóch etapach z wykorzystaniem (po drodze) poziomu związanego z centrum. Rekombinacja taka może być w pewnych przypadkach promienista, zdarza się to jednak bardzo rzadko.

Należy zwrócić uwagę, że omówione modele energetyczne oraz zjawisko rekombinacji dotyczyły półprzewodników zawierających małą ilość domieszek. W przypadku gdy ilość ta jest znaczna, np. powyżej 10^{18} atom/cm³, zniekształcenia pasm, o których wspomniano już w p. 2.1, mogą prowadzić do zatarcia różnic między minimum E_{c1} i E_{c2} w wyniku czego półprzewodnik pośredni może nawet stać się bezpośrednim. Mechanizm rekombinacji zachodzącej w takich półprzewodnikach będzie niewątpliwie odbiegał od wyżej opisanego.

Emisji promieniowania związanej ze zjawiskiem rekombinacji towarzyszy oczywiście absorpcja, którą na rys. 2.4 zaznaczono strzałkami skierowanymi ku górze. Proporcję między tymi procesami można ocenić na podstawie znajomości położenia poziomów quasi-Fermiego. W tym celu, niżej przytoczono rozumowanie prowadzące do określenia warunku, od którego spełnienia zależy możliwość uzyskania wzmocnienia promieniowania w półprzewodniku bezpośrednim [2.2].

Kwant promieniowania może być zaabsorbowany tylko wówczas, gdy w pasmie walencyjnym jest do dyspozycji elektron, który może go zaabsorbować oraz jednocześnie w pasmie przewodnictwa jest nieobsadzony poziom, na który ten elektron może się przenieść. Prawdopodobieństwo takiego zdarzenia jest proporcjonalne do iloczynu

$$f_{nc}(1 - f_{nc}) \quad (2.9)$$

zy czym f_{nv} i f_{nc} — funkcje Fermiego dla elektronów w pasmach walencyjnym i przewodniczym.

Emisja kwantu promieniowania może wystąpić tylko wówczas, gdy pasmie przewodniczym znajduje się elektron, a w pasmie walencyjnym dziura. Prawdopodobieństwo takiej sytuacji jest proporcjonalne do

$$f_{nc}(1 - f_{nv}) \quad (2.10)$$

Tak więc liczba kwantów r_{ab} o energii $h\nu$, absorbowanych na jednostkę czasu, będzie równa

$$r_{ab} = A P_{vc} f_{nc}(1 - f_{nv}) g(\nu) \quad (2.11)$$

gdzie A — liczba kwantów r_{em} o energii $h\nu$, emitowanych w jednostce czasu, będzie równa

$$r_{em} = A P_{vc} f_{nc}(1 - f_{nv}) g(\nu) \quad (2.12)$$

czyli:

$P_{vc} = P_{cv}$ — prawdopodobieństwo wystąpienia przejścia w jednostce czasu, między stanem w pasmach przewodniczym i walencyjnym oraz przeciwnie, jeśli różnica energii między stanami wynosi $h\nu$;

A — pewna stała zawierająca gęstość stanów w pasmach przewodniczym i walencyjnym;

$g(\nu)$ — gęstość fotonów o energii $h\nu$.

Emisja wymuszona przewyższy absorpcję, jeśli będzie spełniona nierówność

$$r_{em} > r_{ab} \quad (2.13)$$

Podstawiając do tej nierówności funkcję Fermiego (2.5) i analogiczną do niej funkcję f_{nv} , otrzymuje się

$$F_c - F_v > h\nu = E_n - E_m \quad (2.14)$$

czyli E_n, E_m — energia poziomów biorących udział w przejściach. Fizyczne uzasadnienie tego kryterium wynika z rys. 2.5, na którym przedstawiono uproszczony model pasmowy z zaznaczeniem poziomów zajętych przez elektrony i dziury. Przejścia promieniste występują wówczas, gdy elektrony z zajętych stanów w pasmie przewodniczym spadają na niezajęte stany w pasmie walencyjnym. Kwant o energii mniejszej od energii odpowiadającej odległości między poziomami quasi-Fermiego może zatem wymusić emisję innego identycznego kwantu przez zajęcie elektronu w doł. Nic może on być natomiast zaabsorbowany, gdyż dla jego energii poziomy w pasmie walencyjnym są nie obsadzone, a pasmie przewodniczym są natomiast zajęte. Dokładnie przeciwnie przedstawia się sytuacja dla kwantów o energii większej od energii odpowiadającej odległości między poziomami quasi-Fermiego. Prawdopodobieństwo wystąpienia emisji wymuszonej jest zatem różne od zera

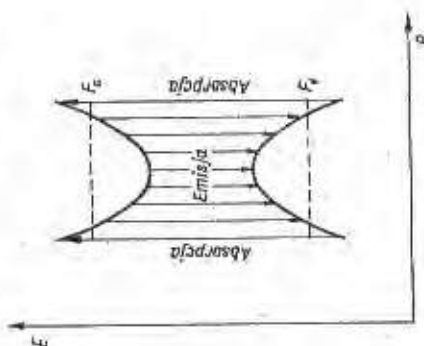
jedynie w przedziale energii $(F_c - F_v) \div E_g$, przy czym zbrocze ilustrującej go krzywej między E_g a maksimum odwzajemnia rozkład gęstości stanów w pasmach przewodniczym i walencyjnym.

Określenie położenia poziomów quasi-Fermiego w wybranych punktach półprzewodnika jest w praktyce zwykle bardzo trudne ze względu na nieznaną lokalną koncentrację wstrzykniętych nośników oraz zniekształcenia pasma występujące przy silnym domieszkowaniu. Wyrażenie (2.14) może być jednak wykorzystane do orientacyjnej oceny koncentracji domieszek w materiale wyjściowym, niezbędnej do uzyskania akcji laserowej.

2.3. Elektroluminescencja w złączach p-n

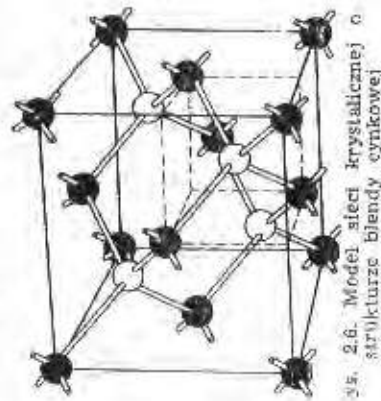
Jednym ze sposobów wywołania zjawiska elektroluminescencji w określonym obszarze półprzewodnika jest wstrzyknięcie do tego obszaru nośników mniejszościowych za pomocą odpowiednio usytuowanego złącza p-n. Zjawisko to zostanie omówione na przykładzie złącza wykonanego z arsenku galu. Wybór tego materiału jest podyktowany z jednej strony stosunkowo dokładnym poznaniem zjawiska elektroluminescencji w złączach p-n wytworzonych z arsenku galu, z drugiej zaś strony tym, że arsenek galu jest dotychczas najpowszechniej stosowanym materiałem do wytwarzania laserów złączonych.

2.3.1. Fizyczne własności arsenku galu. Arsenek galu jest związkiem międzymetalicznym wytworzonym przez syntezę odpowiednich ilości arsenu i galu. Jego kryształy mają sieć krystaliczną o strukturze blendy cynkowej, której model przedstawiono na rys. 2.6. Jedną z cech charakterystycznych sieci o takiej strukturze, w odróżnieniu od sieci typu diamentu (także germanu i krzemu), jest brak elementu symetrii zwanego płaszczyzną inwersji i związana z tym różnica między budową a własnościami kryształu w kierunku [111] oraz w kierunku do niego przeciwnym [111]. Różnica ta występuje również w budowie płaszczyzn (111) oraz (111), co zostało uwidocznione na schematycznym przekroju sieci arsenu galu, wykonanym w płaszczyźnie prostopadłej do płaszczyzny (111) i równoległej do płaszczyzny (110) (rys. 2.7). Powierzchnia ograniczająca kryształ od strony płaszczyzny (111) składa się z atomów galu, podczas gdy powierzchnia do niej przeciwna, będąca płaszczyzną (111), jest

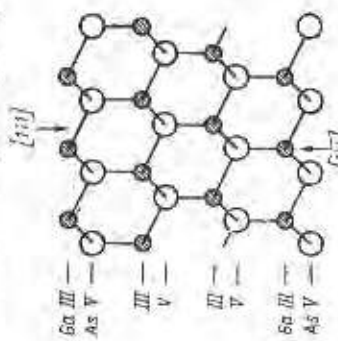


Rys. 2.5. Uproszczony model pasmowy półprzewodnika. Liniami przerywanymi zaznaczono obszar poziomów zajętych przez elektrony i dziury

zbudowana z atomów arsenu. Ponieważ każdy z atomów tworzących te powierzchnie jest związany z atomami sąsiednimi za pomocą trzech wiązań, zatem wszystkie trzy elektrony walencyjne atomów galu są wykorzystane do budowy wiązań, w przypadku arsenu natomiast dwa spośród jego pięciu elektronów walencyjnych pozostają bez przydziału. Konsekwencją tego są znaczne różnice we własnościach płaszczyzn (111)



Rys. 2.6. Model sieci krystalicznej o strukturze blendy cynkowej



Rys. 2.7. Przekrój modelu sieci arsenku galu, wykonany w płaszczyźnie równoległej do płaszczyzn (110) i prostopadłej do płaszczyzn (111)

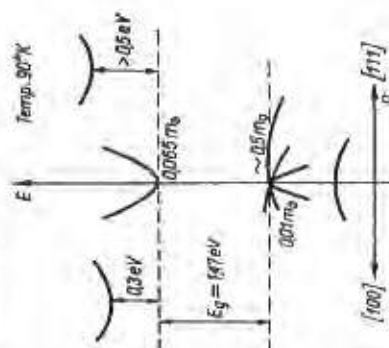
(111). Najważniejsze z nich to szybsze trawienie płaszczyzn (111) w mieszaninach trawiących, łatwiejsze jej rozpuszczanie w ciekłych metalach oraz szybszy wzrost kryształu w kierunku [111]. Różnice te umożliwiają łatwe zidentyfikowanie płaszczyzn (111) i (111). Najczęściej stosowanym w praktyce sposobem jest chemiczne trawienie uprzednio zotworzonego kryształu w odpowiednio dobranej mieszance (np. 75 ml HNO_3 , 15 ml HF , 15 ml CH_3COOH , 0,5 ml Br). W wyniku tego procesu w płaszczyźnie (111) powstają wyraźne jamy trawienia w miejscu, w którym jest ona przecinana przez dyslokację. Jednocześnie płaszczyzna (111) trawi się szybko i równomiernie, dzięki czemu pozostaje gładka błyszcząca.

Opisana asymetria w budowie płaszczyzn (111) i (111) arsenku galu zasadniczo znacznie dla jego własności mechanicznych. W szczególności jest ona przyczyną łatwej łupliwości kryształów GaAs w płaszczyznach równoległych do płaszczyzn (110) i znacznie utrudnionej łupliwości w płaszczyźnie (111). Własność tę można wyjaśnić biorąc pod uwagę rozkład elektrostatycznych sił w kryształach, wynikających z częściowego zjonizowania atomów galu i arsenu. Mianowicie ładunki jonów galu i arsenu mają znak przeciwny, każda z płaszczyzn (111) jest więc nie związana z przylegającą do niej płaszczyzną (111). Każda z płaszczyzn (110) składa się natomiast z takiej samej liczby atomów arsenu i galu, wypadkowa sił wiążących te płaszczyzny jest więc równa zero,

a co za tym idzie można je łatwo od siebie oddzielić. Własność ta jest niezwykle cenna z punktu widzenia technologii laserów złączowych, umożliwia bowiem stosunkowo proste wykonywanie rezonatorów typu Fabry-Perot (rozdz. 3).

Uproszczony model pasmowy arsenku galu przedstawiono na rys. 2.8. Jest rzeczą szczególnie istotną, że obydwa minima pasma przewodnictwa dla pędu $P \neq 0$ są położone znacznie wyżej na skali energii niż minimum centralne. Jak wiadomo z poprzedniego rozdziału, taka budowa pasma przewodnictwa jest cechą półprzewodników bezpośrednich i ma zasadnicze znaczenie dla wydajności procesu elektroluminescencji. Masa efektywna elektronów w pasmie przewodnictwa nie została jeszcze dostatecznie dokładnie określona, przypuszcza się jednak, że wynosi ona ok. $0,065 m_0$, przy czym m_0 jest masą swobodnego elektronu. Pasma podstawowe składa się z dwóch pasm dziur ciężkich, pasma dziur lekkich oraz pasma związanego z oddziaływaniem spinu z orbitą elektronu. Z powodu braku płaszczyzny inwersji w kryształach arsenku galu, wierzchołki pasm dziur ciężkich są względem siebie nieco rozsunięte. Masa efektywna dziur jest w przybliżeniu równa $0,5 m_0$.

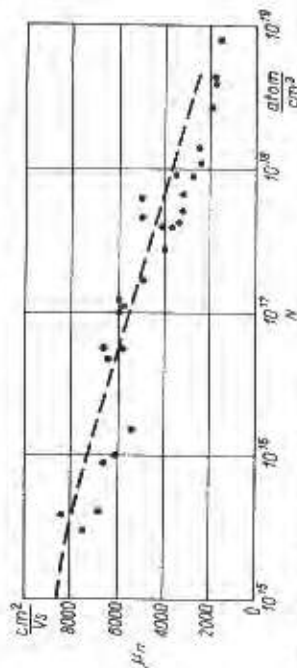
Domieszkami donorowymi dla arsenku galu są niektóre pierwiastki grupy VI. Do najbardziej efektywnych należą selen oraz nieco mniej od niego rozpuszczalny tellur. Pierwiastki te w sieci arsenku galu zajmują pozycje arsenu. Pod względem energetycznym tworzą one pasmo domieszkowe, które dla koncentracji domieszki w kryształach większej od $5 \cdot 10^{17}$ atom/cm³ zlewa się z pasmem przewodnictwa. Rolę akceptorów w arsenku galu odgrywają pierwiastki grupy II, które wchodzą do jego sieci na miejsce atomów galu. Do najefektywniejszych akceptorów należą cynk i kadm. W praktyce najczęściej stosowany jest cynk ze względu na jego dużą rozpuszczalność graniczną (można otrzymać koncentrację ok. 10^{18} atom/cm³) oraz dużą wartość współczynnika dyfuzji. Poziom akceptorowy cynku występuje przy energii o ok. 0,08 eV większej od energii odpowiadającej wierzchołkowi pasma podstawowego. Energia jonizacji tej domieszki szybko maleje w miarę wzrostu jej koncentracji w półprzewodniku. Pierwiastki grupy IV: Si, Ge, Sn, Pb są zwykle donorami, w niektórych przypadkach mogą jednak wchodzić do sieci arsenku galu również jako neutralne atomy. Zachodzi to wówczas, gdy dwa takie atomy zastępują sąsiadujące ze sobą atomy galu i arsenu. Niektóre



Rys. 2.8. Uproszczony model pasmowy arsenku galu. Wartości liczbowe przyporządkowane kropkom odpowiadają masie efektywnej elektronów i dziur

z domieszek, np. tlen, kobalt lub nikiel, tworzą tzw. głębokie poziomy akceptorowe położone w pobliżu środka pasma zabronionego. W przypadku gdy jednocześnie liczba donorów w półprzewodniku jest bardzo niska (np. rzędu 10^{16} atom/cm³), może nastąpić kompensacja donorów przez akceptory, w wyniku czego czystość materiału wzrasta do wartości megamów. Arsenek galu o tak dużej czystości jest nazywany często *półizolującym*.

Ruchliwość elektronów w arsenku galu jest znacznie większa od ruchliwości elektronów w germanie i dochodzi do 8500 cm²/Vs dla bardzo czystych kryształów. Ze wzrostem koncentracji donorów maleje ona według krzywej przedstawionej na rys. 2.9. Ruchliwość dziur jest znacząco

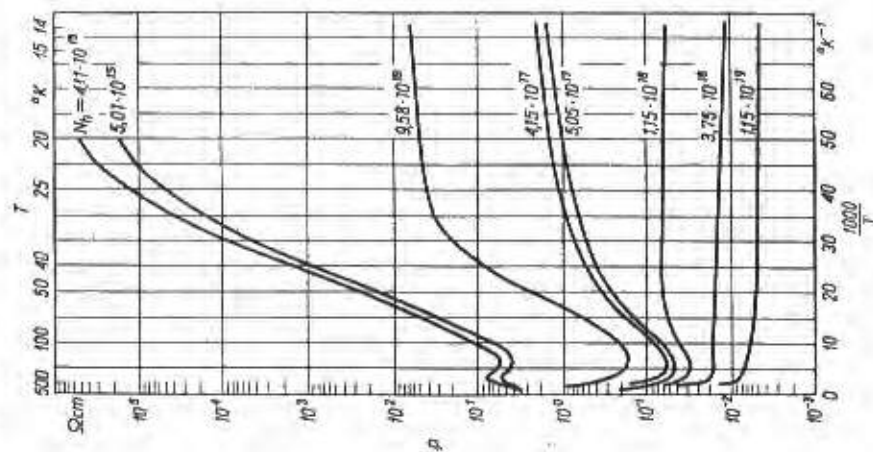


Rys. 2.9. Zależność ruchliwości elektronów od koncentracji domieszek w arsenku galu

nie mniejsza od ruchliwości elektronów i wynosi ok. 400 cm²/Vs dla bardzo czystych próbek oraz odpowiednio mniej dla materiału domieszkowanego. Zarówno ruchliwość elektronów jak i dziur w dużym stopniu zależą od temperatury. Dla ilustracji na rys. 2.11 pokazano taką zależność dla dziur w materiałach zawierających różne ilości tych nośników. Zmiany ruchliwości połączone ze zmianą stopnia jonizacji domieszek w zależności od temperatury pociągają za sobą oczywiście zmianę czystości. Wykres takiej zależności dla materiału typu p domieszkowanego cynkiem pokazano na rys. 2.10.

Ważniejsze własności optyczne i fotoelektryczne arsenku galu będą omawiane przy okazji rozważań zawartych w dalszych rozdziałach. W tym miejscu warto dla porządku wspomnieć, że współczynnik załamania arsenku galu jest stosunkowo duży i wynosi 3,6, co na granicy powietrzem daje współczynnik odbicia $R=0,32$ i kąt graniczny ok. 16°. Własność ta ma istotne znaczenie dla techniki wykonywania rezonatorów (rozdz. 3). Współczynnik absorpcji arsenku galu w dużym stopniu zależy od długości fali, temperatury i koncentracji zawartych w nim domieszek. Liczbowe wartości tego współczynnika podano na rys. 2.12.

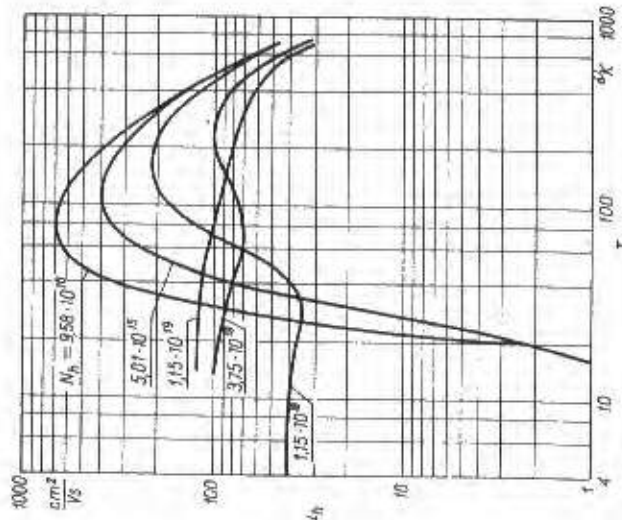
Przewodność cieplna arsenku galu w temperaturze pokojowej jest stosunkowo mała i wynosi ok. 0,7 W/cm °C. Również i ten parametr w dużym stopniu zależy od temperatury i koncentracji domieszek



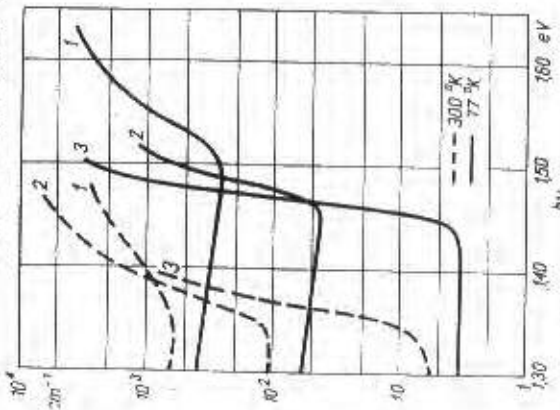
Rys. 2.10. Zmiana rezystywności arsenku galu typu p w zależności od temperatury dla materiałów o różnej koncentracji dziur N_d

(rys. 4.7). Mała przewodność cieplna arsenku galu w zakresie temperatury pokojowej i wyższej od niej stwarza dodatkowe trudności przy wykonywaniu przyrządów z tego materiału.

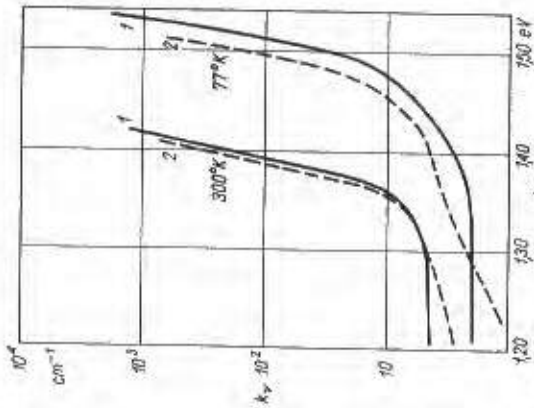
2.3.2. Model pasmowy złącza p-n. Układ pasm energetycznych w złączu p-n bez polaryzacji i po przyłożeniu do niego napięcia w kierunku przewodzenia pokazano na rys. 2.14. Przedstawiony model obrazuje przypadek złącza utworzonego w półprzewodniku zdegenerowanym, gdyż tak właśnie półprzewodniki są z reguły wykorzystywane w technice laserowej z uwagi na konieczność spełnienia warunku (2.14). Jak łatwo zauważyć, przyłożenie odpowiednio dużego napięcia stwarza sytuację, w której elektrony z pasma przewodnictwa po stronie n znajdują



Rys. 2.11. Zależność ruchliwości dziur od temperatury arsenku galu typu p o różnej koncentracji dziur



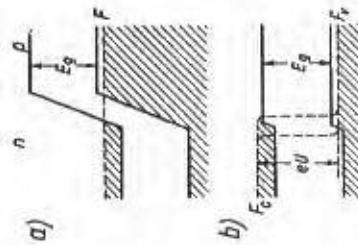
s. 2.12. Współczynnik absorpcji arsenku domieszkowanego cynkiem dla różnych koncentracji akceptorów: 1 — $N_A = 7 \cdot 10^{18} \text{ cm}^{-3}$; 2 — $N_A = 1,6 \cdot 10^{19} \text{ cm}^{-3}$; 3 — $N_A = 1,6 \cdot 10^{20} \text{ cm}^{-3}$ [2.13]



Rys. 2.13. Współczynnik absorpcji arsenku galu domieszkowanego tellurem dla różnych koncentracji donorów: 1 — $N_D = 9,6 \cdot 10^{17} \text{ cm}^{-3}$; 2 — $N_D = 3,0 \cdot 10^{18} \text{ cm}^{-3}$; 3 — $N_D = 1,6 \cdot 10^{19} \text{ cm}^{-3}$ [2.15]

się niejako nad dziurami zajmującymi poziomy pasma walencyjnego po stronie p złącza. Można też zmniejszać odpowiednio położenie poziomów quasi-Fermiego stworzyć sytuację, w której dziury ze strony p będą „przelewały się” na stronę n. Z uwagi na znacznie silniejsze domieszkowanie strony p sytuacja ta wydaje się nawet bardziej prawdopodobna. W rzeczywistości argument większej koncentracji może być prawdziwie dobrane podważony przez znacznie większe uprzywilejowanie wstrzykiwania elektronów do obszaru p niż dziur do obszaru n. Uprzywilejowanie to wynika z większej ruchliwości elektronów i ich mniejszej masy efektywnej, z czym wiąże się głębsze wnikanie poziomu Fermiego w głąb pasma, oraz z krótszego czasu życia elektronów po stronie p z uwagi na większą koncentrację nośników większościowych po tej stronie złącza [2.3].

W sytuacji przedstawionej na rys. 2.14b przez złącze płynie oczywiście prąd elektryczny, którego charakterystykę pokazano na rys. 2.15. Należy zwrócić uwagę, że szczególnie szybko wzrost prądu w miarę



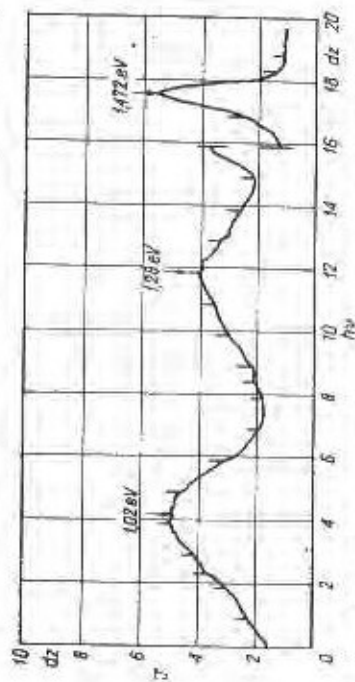
Rys. 2.14. Wykres pasm energetycznych w złączu p-n: a) bez polaryzacji; b) po przyłożeniu napięcia dodatniego $U \approx E_0$

Rys. 2.15. Charakterystyka prądowo-napięciowa diod wykonanych z arsenku galu o różnej koncentracji donorów, dla temperatury $T_s = 4,2^\circ \text{K}$ [2.3]

wzrostu napięcia występuje w pobliżu napięcia ok. 1,5 V, tzn. przy wartości odpowiadającej szerokości pasma zabronionego.

2.3.3. Rekombinacja. Jednoczesna obecność dziur i elektronów w tym samym obszarze półprzewodnika oznacza, jak wiadomo z p. 2.2, możliwość ich rekombinacji. Typowe widmo promieniowania rekombinacyjnego emitowanego przez diodę zasilaną niezbyt dużym prądem (10 A/cm^2) w temperaturze 77°K przedstawiono na rys. 2.16. Dioda ta została wy-

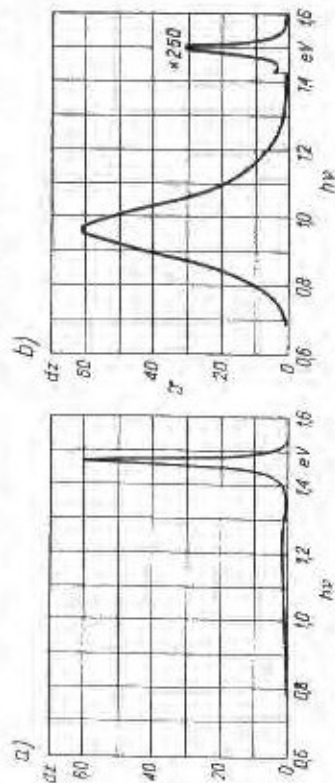
konana za pomocą dyfuzji cynku do arsenku galu domieszkowanego tellurem. W widmie jej promieniowania występują trzy linie emisyjne odpowiadające energiom 1,02, 1,28 i 1,47 eV. Szczegółowe badania wy-



Rys. 2.16. Widmo elektroluminescencji diod GaAs dla temperatury $T_0 = 77^\circ\text{K}$. W prawej części rysunku skala pionowa zmniejszona 25 razy [2.5]

kazały jednak, że kształt widma w rodzaju tego, które przedstawiono na rys. 2.16, zależy ściśle od własności kryształu wyjściowego oraz przebiegu procesu dyfuzji cynku. Szczególnie dużym zmianom podlega część widma od strony małych energii; może się np. zdarzyć, że linia 1,28 eV nie występuje w widmie w ogóle.

Charakterystykę widmową z rys. 2.16 można zinterpretować na podstawie wyników badań nad zjawiskiem fluorescencji. Badania te polegają między innymi na wywołaniu fluorescencji półprzewodnika przez



Rys. 2.17. Widmo fluorescencji próbek GaAs domieszkowanych cynkiem i tellurem dla temperatury $T_0 = 77^\circ\text{K}$: a) GaAs typu p (czarna linia, $N_A = 1,6 \cdot 10^{17} \text{ cm}^{-3}$); b) GaAs typu n (tellur, $N_A = 5,5 \cdot 10^{17} \text{ cm}^{-3}$) (w prawej części rysunku, skala pionowa zmniejszona 250 razy) [3.11]

oświetleniem go światłem o energii większej niż energia pasma zabronionego, a następnie obserwacji emitowanego z powierzchni półprzewodnika światła i pomiarach jego widma za pomocą monochromatora.

Na rys. 2.17 przedstawiono widma fluorescencyjne otrzymane w temperaturze 77°K dla próbek typu n domieszkowanych tellurem i typu p domieszkowanych cynkiem. Pierwiastki te wprowadzają, jak wiadomo, płytkie centra odpowiednio donorowe i akceptorowe.

Emisja z obszaru typu p jest skoncentrowana w wąskiej charakterystyce widmowej (0,02 eV) odpowiadającej energii 1,485 eV — nieco mniejszej od szerokości pasma zabronionego czystego arsenku galu. Występuje ona również przy energii 1,25 eV, jest jednak bardzo mała. Przeciwnie, emisja materiału typu n jest skupiona niemal całkowicie w obszarze małej energii (ok. 0,98 eV) i jedynie bardzo nisko promieniowanie jest obserwowane w pobliżu krawędzi absorpcji (1,51 eV). Ponadto kształt widma emisji materiału typu n jest, w przeciwieństwie do widma dla materiału typu p, niepowtarzalny i zmienia się od kryształu do kryształu.

Porównanie widmowej charakterystyki diody z krzywymi z rys. 2.17 wskazuje na jej znaczne podobieństwo do widmowej charakterystyki fluorescencji materiału typu p. Podobieństwo to sugeruje, że obydwie emisje są wynikiem promienistych przejść tego samego rodzaju, co znajduje zresztą potwierdzenie, jak zobaczymy niżej, w wielu innych faktach doświadczalnych. Kształt charakterystyki emisji diody w zakresie małej energii jest natomiast zwykle bardzo zbliżony do kształtu charakterystyki fluorescencyjnej materiału n, z którego dioda ta została zrobiona. Powyższa obserwacja w połączeniu z faktem, że emisja o małej energii nie występuje w widmie fluorescencyjnym próbek typu p, umożliwia przypuszczenie, że emisja o małej energii jest spowodowana rekombinacją zachodzącą w obszarze typu n. Hipoteza ta jest potwierdzona przez wyniki doświadczeń z bardzo czystym GaAs typu n. Zarówno sam materiał jak i wykonanie z niego diody nie wykazują mianowicie emisji o małej energii.

2.3.4. Zależność widmowej charakterystyki elektroluminescencji od koncentracji domieszek. Z punktu widzenia diod laserowych interesująca jest jedynie fluorescencja występująca w pobliżu krawędzi absorpcji. Położenie wierzchołka charakterystyki tej emisji jest funkcją koncentracji



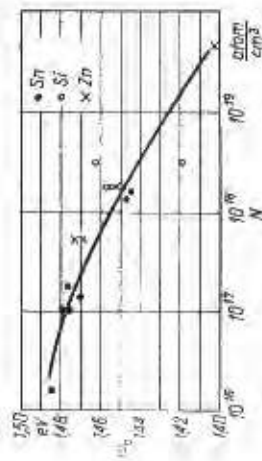
Rys. 2.18. Położenie wierzchołka linii fluorescencyjnej arsenku galu domieszkowanego tellurem i cynkiem jako funkcja koncentracji tych domieszek dla temperatury $T_0 = 77^\circ\text{K}$ [2.11]

tracji domieszek w materiale i zmienia się tak, jak przedstawiono to na rys. 2.18, na którym pokazano krzywe znalezione dla arsenku galu

typu n domieszkowanego tellurem i typu p domieszkowanego cynkiem. Wierzchołek charakterystyki fluorescencyjnej materiału typu n zawierał małe ilości domieszki występuje przy wartościach energii odpowiadających w przybliżeniu szerokości pasma zabronionego. Jej przesunięcie w kierunku większych energii przy wzroście koncentracji domieszek jest prawdopodobnie spowodowane wypełnieniem pasma przewodnictwa przez elektrony zgodnie z modelem opisanym w p. 2.3.6.

Fluorescencja GaAs domieszkowanego cynkiem rozpoczyna się przy energii równej ok. 1,485 eV i ze wzrostem koncentracji przesuwają się w kierunku mniejszych energii. Zjawisko to można interpretować jako rezultat zmniejszenia się szerokości pasma zabronionego E_g . Warto zwrócić uwagę na fakt, że fluorescencja cynku nawet w małych domieszkowanych próbkach występuje przy energiach znacznie mniejszych od szerokości pasma zabronionego dla czystego materiału. Fakt ten sugeruje, że przejście w próbkach typu p kończy się zawsze na poziomach akceptorowych, które przy większych koncentracjach cynku są połączone z pasmem wulencyjnym.

Polożenie wierzchołka charakterystyki elektroluminescencyjnej dla diod wykonanych za pomocą dyfuzji cynku w materiałach typu n o różnej koncentracji donorów oraz diod wykonanych techniką stopową przez wtapianie cynu do GaAs domieszkowanego cynkiem, przedstawiono na rys. 2.19. Zakładając, że koncentracje domieszek N_D i N_A w złączu były



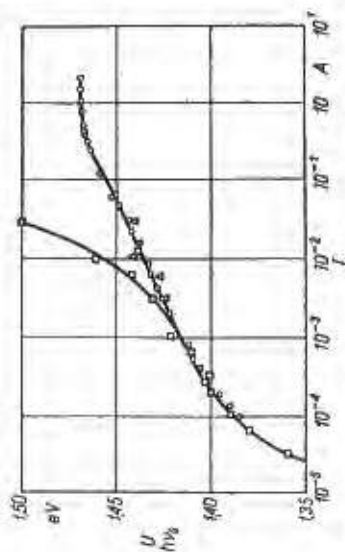
Rys. 2.19. Polożenie wierzchołka widmowej charakterystyki diod GaAs jako funkcja koncentracji domieszek w materiale wyjściowym dla temperatury $T_0 = 77^\circ K$ i gęstości prądu $J = 50 \text{ A/cm}^2$ [2.12]

ówne, co niekoniecznie musi zachodzić z uwagi na złożoną naturę mechanizmu dyfuzji cynku (rozdz. 3), krzywą tę można traktować jako ilustrację związku między koncentracją domieszek a energią emitowanych fotonów. Porównując ją z krzywymi z rys. 2.18 można łatwo zauważyć wyraźne podobieństwo jej kształtu do kształtu krzywej przedstawiającej charakterystykę fluorescencji materiału typu p. Podobieństwem jest zrozumiene w świetle przedstawionej wyżej hipotezy o pochodzeniu emisji diody z obszaru p złącza. Stanowi ono jeden z ważniejszych dowodów, że rekombinacja w diodach obejmuje poziomy akceptorowe. Wznieśniesz zaginięcie się ku dołowi krzywej z rys. 2.19 jest związane prawdopodobnie ze zmniejszaniem się szerokości pasma zabronionego tywołanym obecnością donorów. Przypuszczenie to potwierdzają wyniki

otrzymane dla diod wykonanych techniką stopową z materiału typu p, które zaznaczono na rys. 2.18 krzyżkami.

2.3.5. Zależność widmowej charakterystyki elektroluminescencji od gęstości prądu. Podobnie jak to występowało w przypadku koncentracji domieszek, również i natężenie prądu płynącego przez diodę ma znaczny wpływ na wartość energii odpowiadającej wierzchołkowi charakterystyki elektroluminescencyjnej.

Typową zmianę położenia wierzchołka emisyjnej charakterystyki diody, w zależności od płynącego przez nią prądu wykreślono dla przykładu na rys. 2.20. W środkowym zakresie prądów, w tym przypadku między $2 \cdot 10^{-4} \text{ A}$ a $4 \cdot 10^{-1} \text{ A}$, położenie wierzchołka zmienia się pod wpływem zmian prądu wg zależności typu $I = I_0 \exp\left(\frac{h\nu_0}{\delta}\right)$, przy czym I_0 i δ są pewnymi stałymi współczynnikami. Na krzywą tę dla porównania naniesiono charakterystykę prądowo-napięciową diody. Jeśli uwzględnimy spadek napięcia na szeregowej rezystancji diody, równej ok. 2Ω , otrzymujemy się niemal idealną zbieżność wyników. Oznacza to, że energia,

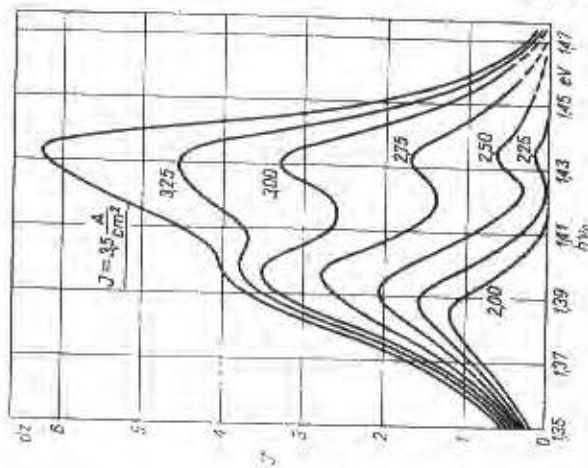


Rys. 2.20. Polożenie wierzchołka widmowej charakterystyki diody w funkcji prądu (kółka), jej charakterystyka prądowo-napięciowa (kwadraty) oraz skorygowana charakterystyka prądowo-napięciowa (trójkąty) [2.18]

przy której występuje maksimum emisji, odpowiada napięciu przyłożonemu do złącza. Przedstawioną korelację między U i $h\nu_0$ otrzymuje się zawsze dla diod wykonanych z materiału o koncentracji donorów $N_D > 3 \cdot 10^{17} \text{ atom/cm}^3$. Uzyskiwane wyniki są niezależne od temperatury, jeśli wprowadzi się poprawkę w skali napięć uwzględniającą zmianę szerokości pasma zabronionego. W diodach, w których koncentracja $N_D < 10^{17} \text{ atom/cm}^3$, obserwowane przesunięcie maksimum widma jest małe. Mało się również zmienia napięcie na diodzie, jak można to stwierdzić na podstawie charakterystyki z rys. 2.15.

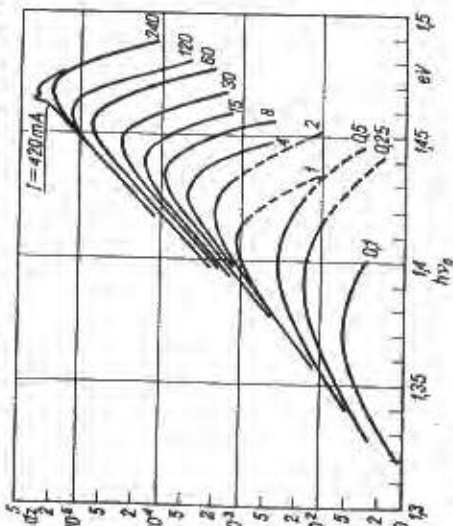
Dokładniejsze badania nad kształtem charakterystyki fluorescencyjnej przy różnych prądach wykazały, że zależność $h\nu_0 = f(I)$ w rodzaju przedstawionej na rys. 2.20 kryje w sobie wiele interesujących zjawisk.

Miedzy innymi stwierdzono występowanie dwóch wierzchołków charakterystyki, których względna wysokość jest funkcją gęstości prądu [2.4]. W zakresie bardzo małych gęstości prądu, 1 do 200 mA/cm² charakterystyka ma jeden wierzchołek, który w miarę wzrostu prądu szybko prze-



Rys. 221. Zmiana kształtu charakterystyki widmowej wywołana prądem dla temperatury $T_0 = 4,2^\circ\text{K}$ [2.4]

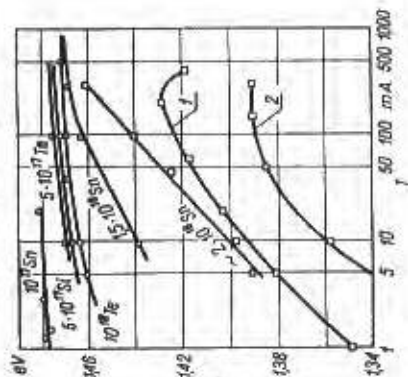
stawa się w kierunku większych energii. W okolicy prądu o gęstości 200 mA/cm² w zakresie nieco większych energii pojawia się drugi wierzchołek. W odróżnieniu od pierwszego wierzchołek ten w miarę wzrostu



Rys. 222. Widmo elektroluminescencji diody GaAs przy różnych prądach zasilania dla temperatury $T_0 = 20^\circ\text{K}$. Dioda wykonana z materiału domieszkowanego selenem do koncentracji $3,2 \cdot 10^{19}$ atom/cm³. Powierzchnia diody $9 \cdot 10^{-4}$ cm² [2.13]

prądu przesuwają się raczej wolno, jego wysokość względna natomiast rośnie znacznie szybciej od wartości prądu (rys. 2.21). Wierzchołek ten zaczyna z czasem dominować nad swym lewostronnym sąsiadem, w wyniku czego przy większych gęstościach prądu, równych ok. 2 A/cm², ponownie obserwuje się tylko jeden wierzchołek przesuwający się w kierunku większych energii. Jego szybkość przesuwania się jest jednak znacznie mniejsza od szybkości, z jaką przesunął się wierzchołek obserwowany przy bardzo małych wartościach prądu. Proces „przesuwania” się tego wierzchołka przedstawiono na rys. 2.22, z którego wynika, że pojęcie przesuwania nie może być tutaj brane dosłownie. Należy raczej uważać, że ze wzrostem prądu zбоче charakterystyki od strony małych energii ulega „nasyceciu”, podczas gdy podnosi się szybko zбоче od strony dużych energii. Jest przy tym rzecz o charakterystyce, że nałożenie promieniowania określone przez to zбоче charakterystyki jest wykładniczą funkcją energii fotonów $h\nu$. Współczynnik w wykładniku tej funkcji jest stały i wynosi 2,5 kT.

Badając zjawisko przesuwania się wierzchołka charakterystyki dla diod wykonanych z materiału o różnej koncentracji domieszki otrzymuje się krzywe przedstawione na rys. 2.23. Krzywe obrazują zachowa-



Rys. 2.23. Zależność położenia wierzchołka charakterystyki widmowej od prądu dla diod wykonanych z materiału o różnej koncentracji domieszki dla temperatury $T_0 = 78^\circ\text{K}$. Złącza dyfundujące cynk, kółkami otrzymano dyfundując $5 \cdot 10^{-19}$ atom/cm³. Złącza epitaksjalne (oznaczone kwadratami) wykonane na podłożu domieszkowanym cynkiem: 1 — $N_A = 3 \cdot 10^{18}$ atom/cm³; 2 — $N_A = 3 \cdot 10^{19}$ atom/cm³ [5.6]

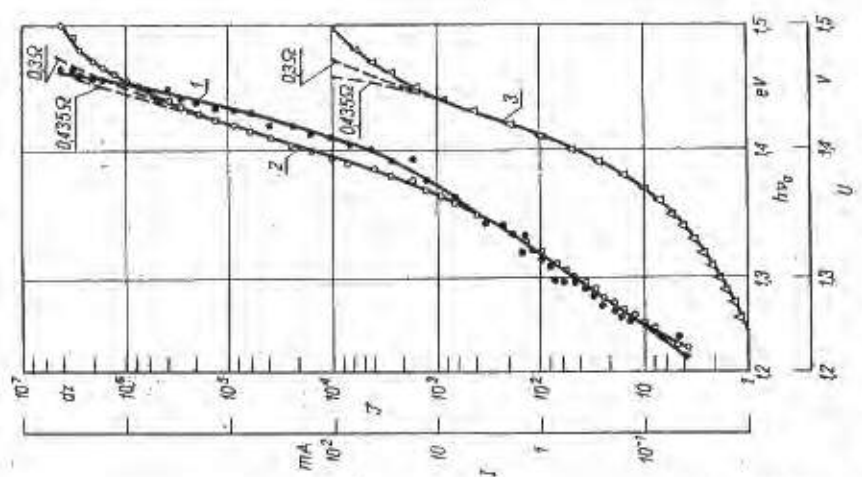
nie się dwóch typów diod: otrzymanych za pomocą techniki dyfuzyjnej i techniki epitaksjalnej. Ich przebieg można wyjaśnić na podstawie opisanego w p. 2.3.6 modelu mechanizmu wypełniania stanów domieszkowych przez wstrzyknięte nośniki.

2.3.6. Mechanizm przejść promienistych i transportu nośników w obszarze złącza. Zarówno mechanizm promienistych przejść nośników jak i procesu ich wstrzykiwania do obszaru, w którym ulegają one rekombinacji, nie został dotychczas dokładnie poznany. Szczególnie duże trudności występują przy ustalaniu poziomów, między którymi zachodzą przejścia. Na podstawie zebranego materiału doświadczalnego można jednak przypuszczać, że przy małych koncentracjach domieszek, a więc

dla $N_D < 10^{17}$ atom/cm³, przejście zachodzi między pasmem przewodnictwa a poziomami akceptorowymi. Dla dużych koncentracji ($N_D > 10^{18}$ atom/cm³) pasma domieszkowe zlewają się z pasmami głównymi i uważa się, że rekombinacja zachodzi między tak zniekształconymi pasmami. Przy jeszcze większych koncentracjach występują znaczne trudności w określeniu stanów biorących udział w przejściach. Można jedynie przypuszczać, że pasma przewodnictwa i walencyjne ulegają bardzo znacznemu zniekształceniu, a nawet, że domieszki powodują zmieszanie się stanów o różnych wartościach wektora falowego k i w związku z tym przestaje obowiązywać reguła wyboru pędu dla przejść bezpośrednich. Zachodzące wówczas przejścia różnią się w sposób zasadniczy od normalnie spotykanych [2.5].

Badania nad zniekształceniem dolnej krawędzi pasma przewodnictwa wywołanego znaczną liczbą atomów domieszek w półprzewodniku (np. ok. 10^{18} atom/cm³) doprowadziły do wniosku, że rozkład stanów w tym pasmie z parabolicznego stałe się wykładniczy, przy czym krawędź pasma znacznie się obniża w kierunku pasma zabronionego tworząc rodzaj „ogona” [2.6]. Wniosek ten znajduje potwierdzenie w wynikach pomiarów charakterystyki widmowej i jej przesunięcia w funkcji prądu (rys. 2.22). Jeżeli bowiem przyjmie się, że stany o wykładniczym rozkładzie energetycznym są wypełniane przez wstrzyknięte nośniki mniejszościowe i że czas życia nośników jest wystarczająco długi, aby mogła ustalić się równowaga termodynamiczna w obszarze wypełnianego pasma, to wierzchołek charakterystyki elektroluminescencyjnej powinien wypadać tuż poniżej energii E_g i przesunąć się ze wzrostem napięcia w sposób obserwowany w doświadczeniach. Ponadto zбоче charakterystyki emisyjnej od strony dużej energii powinno opadać odpowiednio do rozkładu Fermiego, tzn. według funkcji wykładniczej o wykładniku kT (obserwowano ok. 2,5 kT). Zбоче charakterystyki od strony małych energii powinno natomiast odzwierciedlać rozkład gęstości zajętych stanów. Można więc oczekiwać, że na zбочу tej części charakterystyki wystąpi „nasylenie” intensywności emisji jak również, że zбоче to będzie niezależne od temperatury. Wszystkie wymienione zjawiska rzeczywiście są obserwowane w praktyce (rys. 2.22).

Proces transportu nośników w obszarze złącza jest funkcją koncentracji domieszki oraz gęstości prądu płynącego przez złącze. Model mechanizmu transportu, za którym w chwili obecnej przemawia najwięcej faktów doświadczalnych, przynajmniej w zakresie koncentracji domieszek $10^{17} - 2 \cdot 10^{18}$ atom/cm³, opiera się na założeniu, że nośniki są wstrzykiwane w wyniku tunelowania „skośnego” i „poziomego” [2.4]. Koncepcja ta została zaproponowana w wyniku analizy zależności natężenia promieniowania emitowanego przez diodę od przyłożonego do niej napięcia. Przykład takiej zależności przedstawiono na rys. 2.24 jako krzywą zaznaczoną kółkami. Krzywa ta w innym układzie współrzędnych



Rys. 2.24. Natężenie promieniowania w zależności od napięcia na diodzie i energii fotonów odpowiadającej wierzchołkowi linii dla temperatury $T_d = 77^\circ\text{K}$. Poniżej charakterystyka $I-U$; liniami przetykanymi zaznaczono korektę uwzględniającą rezystancję szeregową diody [2.14]; 1 — $I = f(U)$; 2 — $I = f(h\nu_0)$; 3 — $I = f(U)$

na dwa obszary o wyraźnie różnym nachyleniu. W obszarze A zawartym między 1.23 i 1.37 V natężenie promieniowania jest określone zależnością (2.15)

$$I = I_0 \exp(S_A U_A)$$

* w której:

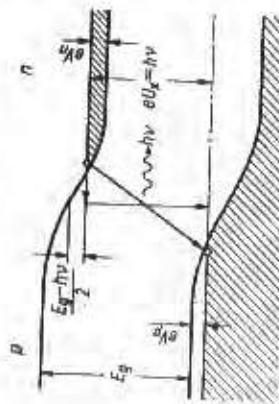
U_A — przyłożone napięcie;

$S_A = 42 \text{ V}^{-1}$;

I_0 — stały współczynnik.

W obszarze B, odpowiadającym $U_A > 1.37 \text{ V}$, nachylenie krzywej rośnie do wartości $S_B = 100 \text{ V}^{-1}$. Jest rzeczą niezwykle interesującą, że wartości S_A i S_B nie zależą od temperatury i są stałe, jak stwierdzono na podstawie pomiarów, aż do temperatury 4.2°K . Zgodnie z klasyczną

teorią przenoszenia nośników przez złącze na zasadzie dyfuzji, nachylenie S powinno być proporcjonalne do T^{-1} . Sprzeczność z doświadczeniem może być wytłumaczona jedynie przyjęciem hipotezy, że ruch nośników odbywa się na zasadzie efektu tunelowego obserwowanego np. w diodach tunelowych [2.7]. Wzrost wielkości nachylenia przy przejściu z obszaru A do B sugeruje, że mamy do czynienia z dwoma odmiennymi mechanizmami tunelowania. Zachowanie się diody w obszarze A można wytłumaczyć na podstawie modelu, który przedstawiono na rys. 2.25. Elektron ze strony n i dziura ze strony p złącza tunelują z poziomów Fermiego ku środkowi obszaru złącza i tam rekombinują emitując foton o energii $h\nu$. Biorąc pod uwagę kierunek ruchu nośników na wykresie z rys. 2.25 tunelowanie to nazwano *skośnym*. Stosownie mechanizmu tunelowania *skośnego* została potwierdzona przez obliczenia prawdopodobieństwa tunelowania elektronu i dziury. Prawdopodobieństwo to dla każdego z poszczególnych aktów przejścia elektronu bądź dziury zostało określone przez Keldysza w następującej postaci [2.8]:



Rys. 2.25. Model mechanizmu „skośnego tunelowania” elektronów [2.4]

przy czym:

$$P_x = \exp(-a \Phi_x \frac{1}{K}) \quad (2.16)$$

$$a = \frac{\pi \sqrt{m^*}}{2eh}$$

m^* — masa tunelującego elektronu (dziury);
 K — natężenie pola elektrycznego w złączu;
 Φ_x — bariera, poprzez którą tunelują elektrony i dziury.

Zakładając w pierwszym przybliżeniu, że

$$V_n + V_p \ll \frac{E_g}{e} - U_x \quad (2.17)$$

można na podstawie rys. 2.25 opisać barierę Φ_x za pomocą wyrażenia

$$\Phi_x \approx \frac{1}{2} (E_g - h\nu) \approx \frac{1}{2} e(V_i - U_x) \quad (2.18)$$

przy czym:

$$E_g \text{ — szerokość pasma zabronionego;}$$

V_n i V_p — potencjał dyfuzyjny;

V_n i V_p zdefiniowano na rys. 2.25.

Łatwo stwierdzić, że

$$V_i = \frac{E_g}{e} + (V_p + V_n) \quad (2.19)$$

Przyjmując teraz, że rozkład domieszek w złączu jest liniowy oraz że o prawdopodobieństwie tunelowania decyduje maksymalne natężenie pola K_m w złączu, wyrażone przez

$$K_m \approx \frac{1.5 (V_i - U_x)}{W} \quad (2.20)$$

otrzymuje się na podstawie wyrażenia (2.16) następujące wyrażenie na natężenie promieniowania:

$$\mathcal{Q} + P_x^2 = \exp(-a W_1 e^{\frac{1}{2}} 1.5 \sqrt{2}) (V_i - U_x)^{\frac{3}{2}} \quad (2.21)$$

przy czym:

W — szerokość złącza;

$$W_1 = W(V_i - U_x)^{\frac{1}{2}} \quad (2.22)$$

Wartość m^* wchodząca do a wynosi $\frac{1}{m^*} = \frac{1}{m_e^*} + \frac{1}{m_h^*}$, przy czym m_e^*

jest masą efektywną tunelującego elektronu, m_h^* — masą efektywną tunelującej dziury.

Nachylenie krzywej $\lg \mathcal{Q} = f(U_x)$ obliczone na podstawie wzoru (2.21) zgadza się bardzo dobrze z wartością znaną na podstawie krzywej z rys. 2.24, co świadczy o słuszności przedstawionego modelu. Ponadto przemawiają za nią inne obserwacje, a mianowicie:

— nachylenie S mierzone dla diod z różnym rozkładem domieszek w złączu jest proporcjonalne do W_1 zgodnie z zależnością teoretyczną;

— częstotliwość odpowiadająca wierzchołkowi widma emitowanego promieniowania odpowiada wartości $\frac{eU_x}{h}$, w której h jest stałą Plancka.

Ten ostatni związek stanie się oczywisty, jeśli weźmie się pod uwagę, że elektrony i dziury znajdujące się bardzo blisko poziomów quasi-Fermiego mają największe prawdopodobieństwo tunelowania, ponieważ bariera potencjału, przez którą powinny tunelować, jest dla nich najmniejsza.

Zmiana nachylenia S dla napięć $U_x > 1.37$ V może być wytłumaczona na podstawie mechanizmu tunelowania poziomego, w którym elektrony ze strony n przechodzą na poziomy pasma przewodnictwa po stronie p . Poziomy ten stanowią omówiony wyżej „ogon” pasma przewodnictwa i ich gęstość zmienia się wykładniczo w miarę zmian energii wg zależności

$$dN(E') = \text{const} \exp(a E') dE' \quad (2.22)$$

przy czym:

E' — energia zdefiniowana na rys. 2.26;

a — pewna charakterystyczna wartość energii dla „ogona” pasma.

Dokładna natura przejścia nazwanego tutaj tunelowym nie jest do tychczas dokładnie znana. Uważa się, że może to być rzeczywiste tunelowanie poprzez barierę potencjału lub też dyfuzja nośników wzdluż

pasem domieszkowych [2.4]. W każdym jednak przypadku jest to przejście bcz zmiany energii Z tego ostatniego faktu wpływa wniosek, że stany w „ogonie” zostaną zapełnione do poziomu E' odpowiadającego poziomowi quasi-Fermiego E' odpowiednio po stronie n złącza. Wartość energii E' można wyznaczyć za pomocą zależności

$$E' = e(V_d - V_a + V_n + U_x) = eE'_{\pi} + e(V_d + V_n - V_i) \quad (2.23)$$

przy czym V_d — potencjał zdefiniowany na rys. 2.26.

Jeśli cykl przepływu prądu jest zakończony przez promieniowanie rekombinacyjne, któremu odpowiada czas życia nośników τ niezależny od energii E' , to natężenie promieniowania $\mathcal{J}$ będzie wielkością proporcjonalną do gęstości stanów $N(E')$, a więc

$$\mathcal{J} = \exp(aE') \quad (2.24)$$

Biorąc pod uwagę zależność (2.23) widzimy, że według wyrażenia (2.24) natężenie jest wykładniczą funkcją napięcia U_x przyłożonego do diody, zgodnie z wynikiem doświadczalnym przedstawionym na rys. 2.24. A tym samym rysunku podano dla porównania zmierzoną zależność między natężeniem promieniowania a energią $h\nu_0$, odpowiadającą wiązce charakterystyki. Charakterystyka ta ma takie samo nachylenie jak krzywa $\mathcal{J} = f(U_x)$, jest jednak względem niej przesunięta o pewną wielką wartość U_0 . Dowodzi to, że wspomniana wyżej odpowiedniość między eU_x a $h\nu_0$ wyraża się następującą zależnością

$$h\nu_0 = e(U_x - U_0) \quad (2.25)$$

Napięcie U_0 stanowi małą część napięcia U_x i jest wielkością zależną od temperatury. Sens fizyczny tego napięcia może być wyjaśniony na podstawie szczegółowej analizy mechanizmu „wypełniania pasma” złączającego dyfuzję nośników po dokonaniu tunelowania oraz faktu, rekombinacja zachodzi między dziurami a elektronami, których energia zawiera się w pewnym przedziale o skończonej szerokości.

Na podstawie przedstawionych wyżej dwóch mechanizmów transportu nośników przez złącze $p-n$ można wyjaśnić opisane w p. 2.3.4 zjawisko występowania dwóch wierzchołków charakterystyki elektroluminescencyjnej przesuwających się z różną szybkością w miarę wzrostu prądu. Stwierdzono mianowicie, że pierwszy wierzchołek wiąże się z tunelowaniem skończonym. Zmiana napięcia przyłożonego do złącza pociąga za sobą zmianę odległości między poziomami biorącymi udział w rekombinacji i stąd wynika duża szybkość przesuwania się wierzchołka cha-

rakterystyki odpowiadającej tej rekombinacji. Wierzchołek przesuwający się wolniej wiąże się z tunelowaniem poziomym. Jego przesunięcie wynika z obsadzenia coraz to wyżej położonych stanów w miarę wypełniania „ogona” przez wstrzykiwane nośniki.

Istnieją przypuszczenia, że obok wymienionych dwóch mechanizmów przenoszenia nośników może istnieć również trzeci mało dotychczas poznany. O jego istnieniu świadczy charakterystyka emisyjna, której wierzchołek występuje przy wartości energii niezależnej od napięcia przyłożonego do diody. Zjawisko to występuje jednak przy małych koncentracjach domieszek ($N_D < 10^{17}$ atom/cm³) i dlatego należy sądzić, że nie odgrywa ono roli w złączach, z których są wykonywane lasery.

Reasumując można więc stwierdzić, że wstrzykiwanie nośników w laserach złączowych, w zakresie prądów przekraczających próg pobudzenia, odbywa się na zasadzie poziomego tunelowania elektronów do pasma przewodnictwa po stronie p złącza.

2.4. Warunki wystąpienia akcji laserowej w złączach $p-n$

Nośniki wstrzyknięte na przeciwną stronę złącza ulegają prawie natychmiastowej rekombinacji z napotkanymi tam nośnikami przeciwnego znaku. W procesie tym wytwarzają się kwanty światła całkowicie przypadkowe w czasie i kierunku. Ich energia jest rozłożona na znacznym obszarze widma i zależy między innymi od temperatury złącza, koncentracji domieszek oraz poziomu iniekcji nośników. Ta część promieniowania ma zatem charakter emisji spontanicznej. Większość wytworzonych fotonów startuje w kierunkach prowadzących do szybkiego opuszczenia obszaru czynnego i nie ma wpływu na zachodzące w nim zjawiska. Obok nich są jednakże i takie, które poruszają się niemal dookoła w płaszczyźnie złącza, dzięki czemu przez znaczny okres czasu pozostają w obszarze z inwersją obsadzeń, oczywiście jeśli nie zostaną wcześniej zaabsorbowane. Fotony te zderzając się z pobudzonymi elektronami mogą wywołać opisany w rozdz. 1 proces emisji wymuszonej. Intensywność tego procesu nasila się szybko w miarę wzrostu gęstości wstrzykniętych nośników. Warunkiem do tego, aby emisja wymuszona skompensowała zachodzące jednocześnie zjawisko absorpcji, jest spełnienie nierówności (2.14). Uzyskuje się wówczas wzmocnienie promieniowania. Najbardziej efektywne w rozwijaniu procesu emisji wymuszonej są oczywiście kwanty o energii odpowiadającej częstotliwości, dla której występuje maksimum wzmocnienia. Promieniowanie o tej częstotliwości lub do niej zbliżonej staje się zatem dominujące w miarę wzrostu natężenia prądu. Pociąga to za sobą ogólne zwięźlenie widmowej charakterystyki emisji oraz szybszy wzrost natężenia promieniowania o wybranej częstotliwości. W celu uzyskania akcji laserowej niezbędne jest jednak, jak wiadomo z rozdz. 1, nadanie układowi wzmacniającemu własności

rezonansowych oraz wprowadzenie dodatniego sprzężenia zwrotnego. Z uwagi na wielkie wzmocnienie optyczne uzyskiwane w złączach p - n sprzężenie to nie musi być silne i może być łatwo zrealizowane przez odbicie części promieniowania od przeciwnie połączonych ścian ograniczających złącze. Z reguły wystarczające jest odbicie na granicy półprzewodnik-powietrze, odbijające płaszczyzny muszą być jednak prostopadłe do płaszczyzny złącza i usytuowane względem siebie pod odpowiednim kątem. Złącze spełniające ten warunek będzie dalej nazywane **rezonatorem**. Rezonator narzuca dalsze warunki na tę część promieniowania spontanicznego, która będzie brała udział w procesie emisji wymuszonej. Mianowicie związane z tym procesem fotony powinny poruszać się nie tylko w płaszczyźnie niezbyt odchylonej od płaszczyzny złącza, lecz również w odpowiednim kierunku względem płaszczyzn odbijających. Prąd, przy którym wzmocnienie układu osiąga wartość wystarczającą do tego, aby liczba fotonów powstających w wyniku emisji była dostateczna do wzbudzenia się drgań w rezonatorze, nosi nazwę **prądu progowego**.

2.4.1. Propagacja promieniowania w rezonatorze. Opisana wyżej emisja fotonów jest oczywiście równoznaczna z generacją fal elektromagnetycznych. Rozpatrując rezonator jako ośrodek, w którym te fale się rozchodzą, można przyjąć, że półprzewodnik z wytworzonym w nim złączem jest obszarem o własnościach dielektryka, którego stała dielektryczna ϵ zmienia się odpowiednio do zmian ładunku przestrzennego i może być zapisana w następującej postaci [2.26]:

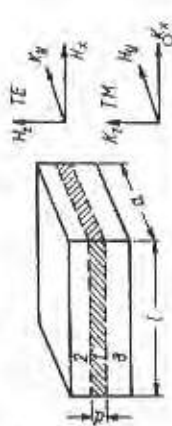
$$\begin{aligned} \epsilon_1 &= \epsilon_1' - \frac{\sigma_1}{j\omega} \\ \epsilon_2 &= \epsilon_2' + \frac{\sigma_2}{j\omega} \\ \epsilon_3 &= \epsilon_3' + \frac{\sigma_3}{jm} \end{aligned} \quad (2.26)$$

przy czym:

$\epsilon_1, \epsilon_2, \epsilon_3$ — składowe rzeczywiste stałe dielektrycznej w obszarach 1, 2 i 3 (rys. 2.27);
 $\sigma_1, \sigma_2, \sigma_3$ — przewodności tych obszarów;
 ω — pulsacja.

Wartości składowych rzeczywistych stałej dielektrycznej są różne dla poszczególnych obszarów rozpatrywanego dielektryka. W szczególności uważa się, że obecność ładunku swobodnych nośników w obszarach p i n przylegających do złącza oraz ich brak w jego obszarze centralnym powoduje nieciągłość wartości ϵ' przy granicach złącza [2.27]. Zagadnienie propagacji promieniowania w rezonatorze sprowadza się zatem do analizy rozchodzenia się fal elektromagnetycznych w warstwie dielektryka ograniczonej przez nieciągłość wartości stałej dielek-

trycznej. Zgodnie z wynikami tej analizy, w warstwie o takich własnościach mogą występować fale rozchodzące się w jej płaszczyźnie i zanikające wykładniczo w kierunku do niej prostopadłym. Rodzaje tych fal, lub ściślej rodzaje rozchodzenia się fal, zwane też często **modami**, mogą być podzielone, podobnie jak to zachodzi w falowodach klasycznych, na mody TE odpowiadające falam elektrycznym poprzecznym oraz mody TM odpowiadające falam poprzecznym magnetycznym (rys. 2.27). Zjawisko ograniczenia warstwy dielektryka płaszczyznami równoległymi do płaszczyzny złącza, wywołane nieciągłością stałej dielektrycznej, powoduje znaczne zmniejszenie strat promieniowania w obszarach 2 i 3 przylegających do obszaru czynnego 1 [2.27, 2.28].



Rys. 2.27. Model rezonatora oparty na elektromagnetycznej teorii propagacji promieniowania w laserze

Płaszczyzny graniczne, równoległe do złącza, łącznie z prostopadłymi do nich płaszczyznami odbijającymi tworzą rezonator, który w najprostszym przypadku ma kształt prostopadłościanu o wymiarach $l \times a \times d$ (rys. 2.27). Wykorzystując w odniesieniu do tego rezonatora wyniki odpowiednio uproszczonej teorii rezonatorów metalicznych można znaleźć wyrażenie określające częstotliwości rezonansowe drgań, które w takim rezonatorze mogą się wzbudzić [1.1]. Wyrażenie to ma następującą postać:

$$\left(\frac{m\pi}{l} \right)^2 + \left(\frac{p\pi}{a} \right)^2 + \left(\frac{q\pi}{d} \right)^2 = \left(\frac{2\pi\nu n}{c} \right)^2 \quad (2.27)$$

przy czym m, p, q — liczby całkowite.

Każdemu układowi tych liczb odpowiada określona długość fali i jej sposób rozchodzenia się w rezonatorze, a więc charakteryzują one mody rezonatora. Interpretacja wyrażenia (2.27) jest bardzo prosta w przypadku szczególnym, gdy drgania w rezonatorze mają charakter fal płaskich rozchodzących się wzdłuż osi podłużnej rezonatora. Wówczas zachodzi warunek $p = q = 0$ i otrzymuje się:

$$\frac{m}{l} = \frac{2\nu n}{c} \quad (2.28)$$

lub

$$\frac{m\lambda}{n} = 2l$$

Częstotliwość rezonansowa odpowiada zatem takiej długości fali, dla której między płaszczyznami odbijającymi ułoży się całkowita liczba m połówek fali, a więc powstanie fala stojąca. Każda liczba m charakteryzuje przy tym określony mod rezonatora. Grupa modów spełniających

zależność (2.28) jest nazywana *modami osiowymi*. W odróżnieniu od nich mody, dla których $p \neq q \neq 0$, noszą nazwę *modów nieosiowych*.

Odległość w sensie długości fali między kolejnymi dozwolonymi modami osiowymi jest określona różnicą $d\lambda$ między długościami fal odpowiadającymi m i $m+1$. Można obliczyć, że odległość ta wyraża się następującym wzorem:

$$\frac{d\lambda}{\lambda^2} = \frac{1}{2n \left(1 - \frac{\lambda}{n} \frac{dn}{d\lambda} \right)} \quad (2.29)$$

2.4.2. Prąd progowy lasera. Fale elektromagnetyczne rozchodzące się w obszarze rezonatora są wzmacniane lub tłumione, co zależy od znaku wypadkowego współczynnika wzmocnienia będącego różnicą współczynnika wzmocnienia α związanego z emisją stymulowaną i współczynnika β określającego straty objętościowe w obszarze czynnym i otaczającym go materiale.

Współczynnik wzmocnienia α jest oczywiście równy co do wartości bezwzględnej współczynnikowi absorpcji k , (rozdz. 1), jego wartość może być zatem znaleziona na podstawie wyrażenia (1.8). W tym celu przyjęto założenie, że charakterystykę absorpcyjną z dostatecznie dobrą dokładnością można opisać za pomocą unormowanej funkcji Lorentza

$$\gamma(\nu) = \frac{\Delta\nu}{2\pi} \frac{1}{(\nu - \nu_0)^2 + \left(\frac{\Delta\nu}{2}\right)^2} \quad (2.30)$$

przy czym:

ν_0 — częstotliwość odpowiadająca środkowi charakterystyki;

$\Delta\nu$ — szerokość charakterystyki mierzona między punktami odpowiadającymi połowie mocy.

Dzięki unormowaniu

$$\int_{-\infty}^{+\infty} \gamma(\nu) d\nu = 1 \quad (2.31)$$

współczynnik wzmocnienia $\alpha(\nu)$ może być przedstawiony w następującej postaci:

$$\alpha(\nu) = \int k \cdot d\nu \cdot \gamma(\nu) \quad (2.32)$$

co po podstawieniu wartości całki ze wzoru (1.8) oraz uwzględnieniu wartości A_{nm} daje wyrażenie

$$\alpha(\nu) = \frac{c^3}{8\pi\nu_0^3 n^2 \tau} \frac{g_n}{g_m} \left(N_m - \frac{g_m}{g_n} N_n \right) \gamma(\nu) \quad (2.33)$$

w którym:

N_n i N_m — liczba atomów w 1 cm^3 znajdujących się na poziomach wyższym i niższym;

g_n i g_m — liczby określające krotność poziomów;

n — współczynnik załamania;

τ — czas życia nośnika odpowiadający przejściu spontanicznemu z poziomu wyższego na poziom niższy.

W rozwinięciu akcji laserowej najbardziej efektywne jest promieniowanie o częstotliwości ν_0 , dla której występuje maksimum wzmocnienia. W dalszych rozważaniach będzie przyjęto brana pod uwagę wartość $\alpha(\nu)$ w punkcie ν_0 , którą po dokonaniu odpowiednich podstawień do wzoru (2.33) można wyrazić za pomocą zależności

$$\alpha(\nu_0) = \frac{\lambda^2}{4\pi^2 n^2 \tau} \left[N_m \left(\frac{g_n}{g_m} \right) - N_n \right] \frac{1}{\Delta\nu} \quad (2.34)$$

w której $\Delta\nu = \frac{2}{\pi\gamma(\nu_0)}$ — szerokość charakterystyki emisji spontanicznej.

W przypadku temperatur bliskich 0°K można przyjąć $N_m = 0$, co upraszcza wyrażenie (2.34) do postaci

$$\alpha(\nu_0) = - \frac{\lambda^2}{4\pi^2 n^2 \Delta\nu} \frac{N_n}{\tau} \quad (2.35)$$

Iloraz $\frac{N_n}{\tau}$ określa liczbę kwantów emitowanych w wyniku emisji spontanicznej w ciągu 1 s. Wielkość ta może być określona za pomocą wzoru

$$\frac{N_n}{\tau} = \frac{I\eta_{qw}}{e} \frac{1}{Ad} = \frac{J\eta_{qw}}{e} \frac{1}{d} \quad (2.36)$$

w którym:

I — prąd płynący przez złącze;

e — ładunek elektronu;

η_{qw} — wewnętrzna wydajność kwantowa (p. 5.6);

A — pole powierzchni złącza;

d — grubość obszaru czynnego;

J — gęstość prądu.

Podstawiając tę zależność do wzoru (2.35) otrzymuje się wyrażenie

$$\alpha(\nu_0) = \frac{\lambda^2}{4\pi^2 n^2 \Delta\nu} \frac{J\eta_{qw}}{ed} \quad (2.37)$$

Wzbudzenie lasera może wystąpić tylko wówczas, gdy wypadkowe wzmocnienie fali w granicach obszaru czynnego w ciągu jednego kompletnego przejścia, tj. od jednego zwierciadła do drugiego i z powrotem, jest wystarczające do skompensowania strat wynikających z niezupełnego odbicia promieniowania na końcach rezonatora. Warunek ten może być wyrażony w następującej postaci:

$$R_1 R_2 \exp(\alpha - \beta) 2l = 1 \quad (2.38)$$

przy czym:

R_1, R_2 — współczynniki odbicia zwierciadeł;

l — długość rezonatora.

W przypadku gdy $R_1 = R_2$, wyrażenie (2.38) przybiera postać

$$R \exp(\alpha - \beta) l = 1 \quad (2.39)$$

Biorąc pod uwagę otrzymany wynik oraz wyrażenie (2.37) można znaleźć wartość progowej gęstości prądu, tzn. natężenie prądu na jednostkę powierzchni złącza, przy którym jest spełniony warunek wzburzenia się drgań

$$J_{pr} = \frac{4\pi^2 e n^2 d \Delta \nu}{\eta_{qc} l^2} \left(\frac{\ln \frac{1}{R}}{\beta + \frac{1}{l}} \right) \quad (2.40)$$

W praktycznie stosowanym układzie jednostek wyrażenie na gęstość prądu w $\frac{A}{cm^2}$ przyjmuje postać

$$J_{pr} = \alpha \left(\beta + \frac{1}{l} \right) \left(\frac{\ln \frac{1}{R}}{l} \right) \quad (2.41)$$

przy czym:

$$\alpha = 9,89 \cdot 10^4 \frac{n^2 E^2 d \Delta E}{\eta_{qc} l}, \text{ w } A/cm; \quad (2.42)$$

E — energia odpowiadająca długości promieniowanej fali, w eV;

ΔE — szerokość linii emisji spontanicznej, w eV;

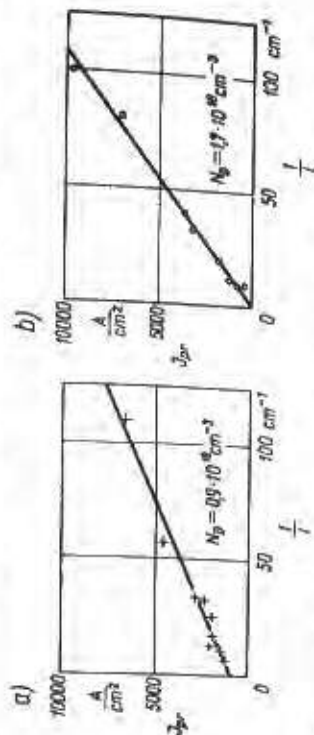
d — grubość obszaru czynnego, w cm;

η_{qc} — wydajność kwantowa;

β — współczynnik strat objętościowych, w cm^{-1} ;

l — długość rezonatora, w cm.

Należy zwrócić uwagę, że współczynnik α w otrzymywanym wyrażeniu jest również pewną funkcją prądu, gdyż zależą od niego wielkości $E_0(l)$, η_{qc} i ΔE . Wyniki doświadczeń wskazują jednak, że współczynnik α może być uważany z dobrym przybliżeniem za wartość stałą. Tak np.



Rys. 2.28. Gęstość prądu progowego wykreślona jako funkcja odwrotności długości rezonatora Fabry-Perot [2.8] dla różnych koncentracji donorów: a) $N_D = 0.9 \cdot 10^{18} \text{ cm}^{-3}$; b) $N_D = 1.7 \cdot 10^{18} \text{ cm}^{-3}$

uzyskana w drodze doświadczalnej zależności między wartością prądu progowego a odwrotnością długości rezonatora, którą przedstawiono na rys. 2.28, wykazuje proporcjonalność tych dwóch wielkości.

Przedstawione na rys. 2.28 wyniki umożliwiają ocenę wartości współczynników α i β , które wynoszą odpowiednio [2.9]:

$$\alpha = 41.7 \text{ A/cm}, \quad \beta = 14 \text{ cm}^{-1} \quad \text{dla } N_D = 0.9 \cdot 10^{18} \text{ atom/cm}^3$$

$$\alpha = 71.3 \text{ A/cm}, \quad \beta = 3 \text{ cm}^{-1} \quad \text{dla } N_D = 1.7 \cdot 10^{18} \text{ atom/cm}^3$$

Wartość parametrów α i β można również znaleźć badając zależność prądu progowego od współczynnika odbicia R . Jak wynika ze wzoru (2.41), między tymi wielkościami zachodzi następujący związek:

$$\frac{1}{\alpha} \ln J_{pr} - \beta l = \ln \frac{1}{R} \quad (2.43)$$

Półlogarytmiczny wykres $\frac{1}{R}$ w funkcji prądu progowego jest zatem linią prostą, z której przebiegu łatwo odczytać obydwa szukane parametry. Pomiar przeprowadzone na laserach pokrywanych warstwą dielektryczną lub srebrem dały następujące wyniki [2.10]:

$$\alpha = 200 \div 24 \text{ A/cm} \quad \text{oraz} \quad \beta = 10 \div 60 \text{ cm}^{-1}$$

Badania te wykazały również, że aproksymacja zależności między prądem progowym a wzmocnieniem za pomocą funkcji liniowej jest dostatecznie dobra.

Przytoczone wyżej liczby określające wartość β wykazują znaczny rozrzut. Poza błędem, jakim niewątpliwie były obarczone pomiary, rozrzut ten można tłumaczyć złożoną naturą przyczyn decydujących o wartości tego parametru. Straty promieniowania oznaczone współczynnikiem β wynikają z absorpcji promieniowania w obszarze czynnym oraz jego dyfrakcji i rozpróśnienia, które prowadzą do absorpcji w części materiału przylegającej do obszaru czynnego. Wartość tych strat jest funkcją wymiarów obszaru czynnego, długości fali i współczynnika załamania oraz współczynnika absorpcji. Absorpcja może być związana z procesem odwrotnym do emisji (rozdz. 1) i nosi wówczas nazwę absorpcji krawędziowej określonej jako przenoszenie elektronów z pasma walen- cyjnego do pasma przewodnictwa kosztem energii fotonów. Udział takiej absorpcji w obrębie obszaru czynnego został uwzględniony przy wyważaniu wzoru (2.33). Obok absorpcji krawędziowej występuje tzw. absorpcja na swobodnych nośnikach, w wyniku której elektrony kosztują energii uzyskanej od fotonów przesuwają się na wyższe poziomy energetyczne w obrębie swego pasma. Spadając z nich następnie na niższe, nie obsadzone poziomy, elektrony oddają nadmiar energii do sieci kryształu, gdzie zostaje ona zamieniona na ciepło.

Liczbowa wartość absorpcji niezależnie od jej mechanizmu określa współczynnik absorpcji k . Jest on funkcją energii fotonów, szerokości pasma zabronionego w półprzewodniku, koncentracji znajdujących się

w nim domieszek oraz temperatury. Wartości współczynnika absorpcji dla arsenku galu typu p i typu n domieszkowanych odpowiednio cynkiem i tellurem dla różnych koncentracji występujących w laserach złączowych przedstawiono na rys. 2.12 i 2.13. Zakres, w którym krzywe mają większe nachylenie, odpowiada absorpcji krawędziowej, początkowe odcinki krzywych określają natomiast absorpcję na swobodnych nośnikach.

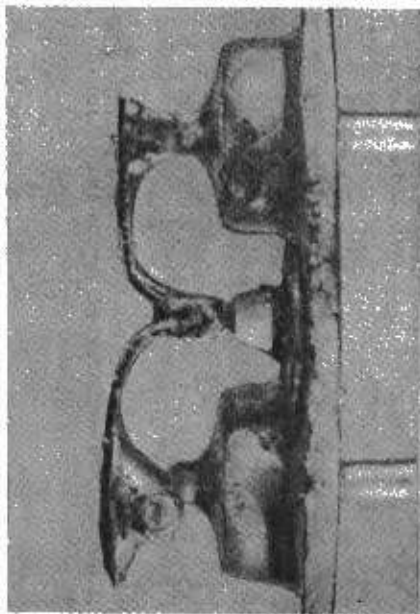
Warto podkreślić, że w laserach spotykanych w praktyce dodatkową przyczyną strat są niewątpliwie wady wnętrza rezonansowej, wskutek których znaczna część promieniowania wychodzi poza obszar czynny nie przyczyniając się do dalszego rozwoju emisji wymuszonej.

3. KONSTRUKCJA I TECHNOLOGIA LASERÓW ZŁĄCZOWYCH Z ARSENKU GALU

3.1. Ogólny opis budowy laserów

Zgodnie z rozważaniami rozdz. 2, laser złączowy jest diodą p-n o odpowiednim kształcie złącza, które pełni jednocześnie funkcję wnętrza rezonansowej. Zadaniem tej wnęki jest stworzenie warunków do wzbudzenia się oscylacji oraz uprzywilejowanie drgań o pożądanych modach.

Klasyczną konstrukcję lasera zmontowanego na oprawce tranzystorowej pokazano na rys. 3.1. Szkic samego rezonatora typu Fabry-Perot



Rys. 3.1. Klasyczna konstrukcja lasera z rezonatorem typu Fabry-Perot zmontowanym na oprawce tranzystorowej [8.16]

przedstawiono na rys. 3.2. Złącze p-n znajduje się zwykle w odległości $50\text{--}20\text{ }\mu$ od powierzchni płytki w płaszczyźnie równoległej do tej powierzchni. Dla promieni biegnących wzdłuż osi podłużnej lasera ściana przednia i tylna rezonatora spełniają funkcję półprzewodzących zwierciadeł. Ściany te są do siebie równoległe a ich powierzchnia bardzo gładka. W celu uniknięcia możliwości powstania akcji laserowej w kierunku poprzecznym, boczne powierzchnie rezonatora są matowe i nie-

Macierz fotodiod

Linijka diodowa jest stosowana w spektrometrii do rejestracji widma optycznego. Linijka typu S3903 (MOS Linear Image Sensor) składa się z 1024 fotodiod mających wspólny wzmacniacz i generator impulsowy. Wszystkie fotodiody wytworzone są na jednym podłożu krzemowym z krzemu typu p (anody o grubości $400\mu\text{m}$) poprzez naniesienie na nie paski z warstwy krzemu typu n (katody o rozmiarach $2.5\text{mm} \times 1\mu\text{m}$). Każdy taki pasek stanowi fotodiode, która ma pojemność $C=10\text{pF}$.

Gdy na taką macierz fotodiod pada światło, to powstaje ładunek $Q = iT$ (i – prąd płynący przez fotodiode proporcjonalny do intensywności światła, T – czas ładowania). Ładunek ten doprowadza do powstania na kondensatorze, jakim jest macierz, napięcia $V = Q/C$. Każdy z kondensatorów jest rozładowywany po kolei, co powoduje powstanie szeregu impulsów, których amplituda odpowiada intensywności światła padającego na odpowiedni element macierzy. Rozładowanie realizuje się ~~na~~ przez jedną szynę łączącą wszystkie elementy macierzy poprzez tranzystory polowe. Generator powoduje „zamknięcie kolejnego klucza” i rozładowanie kolejnego elementu macierzy. Czas rozładowania jednego elementu macierzy jest rzędu ns, odpowiednio rozładowanie całej linijki odbywa się w czasie kilku ms. Częstotliwość powtarzania procesu (odczytywania macierzy) jest rzędu 1kHz. Impulsy z szyny trafiają do wzmacniacza, następnie na przetwornik ADC (analog-digital converter), który przekształca je w odpowiedni szereg zakodowanych cyfr. Z ADC cyfrowy szereg kodowy przekazywany jest do rejestrów pamięci komputera. W efekcie na ekranie monitora (dzięki zastosowaniu odpowiedniego programu) można oglądać widmo.

Parametry opisujące macierz, to:

1. zakres widmowy - dla S3903 (200; 1000)nm
2. czułość maksymalna - dla S3903 przy 650nm
3. zakres dynamiczny (stosunek maksymalnego sygnału do minimalnego sygnału) - dla S3903 jest 3 rzędy