

République Algérienne Démocratique et Populaire

Ministère de l'Enseignement Supérieur et de la Recherche Scientifique

Université Ferhat Abbas de Sétif 1



Polycopié de cours

Module :

physique des semi-conducteurs

Présentée à la faculté des sciences

Département de Physique

Par

REFFAS Mounir

2024

AVANT PROPOS

Ce polycopié est destiné aux étudiants de la première année master, spécialité physique énergétique et énergies renouvelables. Nous pensons que ce cours permettra à l'étudiant de comprendre les contextes scientifiques et technologiques dans le domaine des semi-conducteurs. Le contenu de ce module est organisé en cinq chapitres permettant au lecteur d'apprendre aussi les différentes propriétés physiques des matériaux semi-conducteurs et leurs applications dans le domaine des énergies renouvelables.

Dr. Mounir REFFAS

SOMMAIRE

Chapitre 1. Généralités

1.1. Théorie des bandes d'énergie	1
1.1.1. Les orbitales atomiques à la structure de bandes	1
1.1.1.Orbitales atomiques.....	1
1.1.1.2.Orbitales moléculaires.....	3
1.2.Bandes d'énergie.....	6
1.2.1 Cas du solide : Conducteur - Isolant – Semiconducteur	7
1.3.Semi-conducteurs à l'équilibre thermodynamique	8
1.3.1. Notion de trou.....	8
1.3.2. Notion de gap	9
1.3.2.1.Nature du Gap	10
1.3.2.2. Masse effective et relation de dispersion	12
1.4.Structure des bandes d'énergie	14
1.4.1.Arséniure de gallium	15
1.4.2.Silicium	16
1.5.Densité d'états	16
1.5.1.Densité d'états d'électrons et trous	18

Chapitre 2. Semi-conducteurs Intrinsèques et Extrinsèques

2.1. Un semi-conducteur est dit « intrinsèque » ?	21
2.2. Dopage	22
2.2.1. Exemple de dopage	24
2.2.3. Semi-conducteur de type n.....	26
2.2.4. Semi-conducteur de type p.....	27
2.2.5. Niveaux d'énergie des impuretés.....	28
a) Impuretés donneurs d'électrons	28
b) Impuretés accepteurs d'électrons	28
2.2.6. Dopants amphotères et autodopage.....	29
2.2.7. Niveaux profonds	29
2.2.8. Dopage fort.....	30
2.3. Un semi-conducteur est dit « extrinsèque » ?	30

2.3.1. Pour un semi-conducteur de type n :	31
2.3.2. Pour un semi-conducteur intrinsèque	32

Chapitre 3. Phénomène de Génération-Recombinaison

3.1. Etude des Semi-conducteurs hors équilibre	33
3.1.1. Définition d'un semi-conducteur hors équilibre	33
3.1.2. Calcul des densités de courant	33
3.1.2.1. Calcul de la densité de courant de conduction (mobilité de porteurs de charges)	33
3.1.3. Calcul de la densité de courant de diffusion	34
3.2. Phénomène de Génération-Recombinaison	37
3.2.1. Recombinaison directe (électron – trou)	38
3.2.2. Recombinaison indirecte (centre de recombinaison)	40
3.2.2.1. Cas d'un semi-conducteur de type (N).....	41
3.2.2.2. Cas d'un semi-conducteur de type (P)	41
3.2.3. Zone de déplétion (dépeuplée)	42

Chapitre 4. Etude des Jonctions PN

4.1. Définition des jonctions PN (abrupte et graduelle).....	43
4.1.1. Jonction PN abrupte :	43
4.1.2. Jonction graduelle :	43
4.2. Etude d'une jonction PN abrupte non polarisée à l'équilibre	44
4.2.1. La charge d'espace $\rho(x)$: ($x_P < x < x_N$).....	44
4.2.2. Calcul des champs électriques : $E_P(x)$ et $E_N(x)$	47
4.2.3. Calcul des potentiels électriques : $V_P(x)$ et $V_N(x)$	48
4.2.4. Calcul de la tension de diffusion V_d ; barrière de potentiel.....	50
4.2.5. Calcul de la largeur de la zone de charge d'espace w	52
4.3. Etude d'une jonction PN polarisée (hors équilibre)	53
4.3.1. Polarisation directe.....	53
4.3.2. Polarisation inverse.	54
4.3.2. Caractéristique courant – tension $I(V)$	55
4.4. Types de jonctions PN.....	56

4.4.1. Diodes Zener	56
4.4.2. Diodes à avalanche.....	57
4.4.3. Diode à effet tunnel	57
4.4.5. Photodiode.....	58
4.4.6. Diode électroluminescente	58
4.4.7. Diode Laser	58
4.4.8. Cellules solaires.....	58
4.5. Applications des jonctions PN :	59
4.5.1. Redressement de signaux alternatifs	59
4.5.2. Commutation.....	60

Chapitre 5. Dispositifs à hétérojonction

5.1. Jonction métal-métal	61
5.2. Jonction métal-semiconducteur.....	62
5.2.1. Diode Schottky	62
5.2.2. Jonction Ohmique	64
5.3. Transistors bipolaires	65
5.3.1. Introduction	65
5.3.2. Définitions.....	65
5.3.1. Transistor NPN non polarisé	67
5.3.2. Transistor NPN polarisé.....	67
5.3.2.1. Courants à travers les jonctions.....	69
5.3.2.2. Effet transistor : Gains en courant.....	69
5.3.3. Réseau des caractéristiques statiques du transistor NPN	71
5.3.3.1. Montage émetteur commun.....	71
5.3.4. Régimes de fonctionnement du transistor	74
5.3.5. Limitations physiques	75
5.3.5.1. Tensions de claquage	75
5.3.5.2. Limitation en puissance.....	75
5.3.5.3. Limitation du courant maximum.....	76

Références bibliographiques.....	77
---	-----------

Chapitre 1.

Généralités

1.1. Théorie des bandes d'énergie

1.1.1. Les orbitales atomiques à la structure de bandes

1.1.1.1. Orbitales atomiques

Considérons l'atome d'hydrogène, le plus simple modèle de tous les atomes. Il est constitué d'un électron de charge $-e$ gravitant autour d'un noyau de charge $+e$. L'Hamiltonien de l'électron dans son mouvement autour du noyau s'écrit :

$$H = -\frac{\hbar^2}{2m^*} \nabla^2 + V(r)$$

$-\frac{\hbar^2}{2m^*}$: représente son énergie cinétique, $V(r)$: son énergie potentielle dans le champ coulombien du noyau, ∇^2 l'opérateur est la somme des dérivées secondes partielles, il est aussi noté Δ et appelé Laplacien. L'équation aux valeurs propres s'écrit:

$$H\phi_{nlm}(r, \theta, \phi) = E_{nl} \phi_{nlm}(r, \theta, \phi)$$

n est le nombre quantique principal, l est le nombre quantique azimutal, il représente le moment angulaire de l'électron dans son mouvement autour du noyau. m représente la projection du moment angulaire sur un axe choisi comme axe de quantification. Les valeurs possibles de n sont $n = 1, 2, \dots$, pour chaque valeur de n il existe n valeurs de l , $l = 0, 1, \dots, n-1$, pour chaque valeur de l il existe $2l+1$ valeurs de m ,
 $m = -l, -(l-1), \dots, l-1, l$. Dans la terminologie courante les états correspondant aux différentes valeurs de l portent des noms spécifiques.

$$\text{Etat } s \quad l = 0$$

$$\text{Etat } p \quad l = 1$$

$$\text{Etat } d \quad l = 2$$

$$\text{Etat } f \quad l = 3$$

On écrit les différents états électroniques sous la forme 1s; 2s, 2p; 3s, 3p, 3d;

A chaque fonction propre ϕ_{nlm} correspond une valeur propre E_{nlm} . Dans le cas de l'atome isolé, en raison de la symétrie sphérique du potentiel les niveaux correspondant aux mêmes valeurs de n et de

l sont tous confondus, $E_{nlm} = E_{nl}$ quel que soit m . En outre, dans le cas de l'atome d'hydrogène, en raison de la nature en $1/r$ du potentiel coulombien, les niveaux correspondant aux mêmes valeurs de n sont tous confondus, $E_{nl} = E_{nl}$ quel que soit l . Si le zéro de l'énergie correspond à l'électron complètement séparé du noyau sans énergie cinétique, les niveaux d'énergie E_n sont donnés (pour $n \geq 1$) par :

$$E_n(\text{eV}) = -R \frac{1}{n^2}$$

$$R = \frac{m_e e^4}{2\hbar^2} = 13.6 \text{ eV}$$

R est La constante de Rydberg s'exprime en m^{-1}

$$R = \frac{m_e e^4}{8\varepsilon_0^2 \hbar^3 c} = 1.097373 \times 10^7 \text{ m}^{-1}$$

Les orbitales s ($l = 0$) et p ($l = 1$) sont représentées sur la figure (1.1). Le rayon de Bohr de l'état fondamental.

$$a_0 = \frac{\hbar^2}{m_e e^4} = 0.529 \text{ \AA}$$

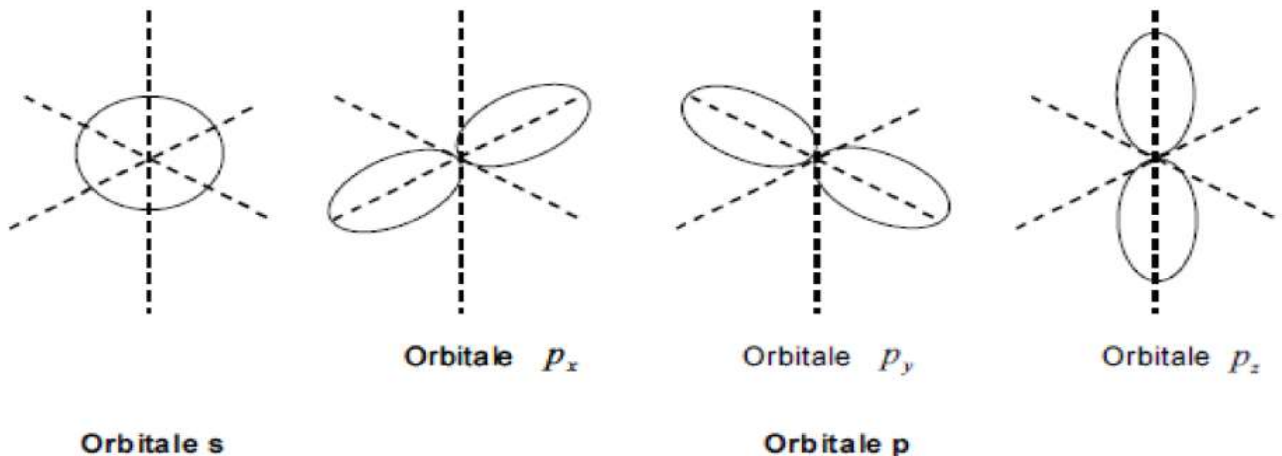


Figure 1.1 : Orbitales atomiques s et p

Dans le cas plus général, un atome de nombre atomique Z est constitué d'un noyau de charge $+Ze$ et d'un nuage électronique de Z électrons. Les électrons de la couche externe sont appelés électrons de

valence, les électrons des couches internes sont appelés électrons du cœur. Les électrons de valence jouent un rôle essentiel d'une part dans les propriétés chimiques des matériaux et d'autre part dans leurs propriétés électriques.

Le potentiel que voit un électron de valence est le potentiel coulombien du noyau écranté, d'une part par les électrons du cœur et d'autre part par les autres électrons de valence. Le potentiel résultant peut encore être supposé central, c'est-à-dire à symétrie sphérique, mais il ne varie plus en $1/r$ comme dans le cas unique de l'atome d'hydrogène. Il en résulte que les états correspondant aux différentes valeurs de m restent dégénérées, mais que les états correspondant aux différentes valeurs de l ont des énergies différentes. La différence d'énergie $np - ns$ augmente à mesure que le potentiel s'éloigne d'une variation en $1/r$, c'est-à-dire à mesure que le nombre d'électrons de valence augmente. En d'autres termes, sur une même ligne du tableau périodique, l'écart d'énergie $np - ns$ augmente avec le nombre atomique. Sur la ligne du Lithium par exemple, les valeurs sont les suivantes (W.A.Harrison, 1980).

	Li	Be	B	C	N	O	F
$E_{2p} - E_{2s}(\text{eV})$		4.04	5.88	8.52	11.52	15.04	18.84

1.1.1.2. Orbitales moléculaires

Considérons la molécule la plus simple d'hydrogène. Elle est constituée de deux noyaux autour desquels gravitent deux électrons. Négligeons l'interaction électron-électron, chacun des électrons est dans un potentiel résultant des potentiels coulombiens de chacun des noyaux. L'Hamiltonien d'un électron s'écrit :

$$H = -\frac{\hbar^2}{2m^*} \nabla^2 + V(r - R_1) + V(r - R_2)$$

où R_1 et R_2 représentent les positions des noyaux. La distance $R_1 R_2$ est de 0,74 Å dans la molécule d'hydrogène. L'équation aux valeurs propres s'écrit :

$$H \psi(r) = E \psi(r)$$

Les fonctions d'onde $\psi(r)$ sont appelées *orbitales moléculaires*.

Quand les deux atomes sont suffisamment éloignés, les états électroniques sont représentés par des orbitales atomiques. Au fur et mesure que les deux atomes se rapprochent, ces états atomiques se couplent pour donner naissance à des états moléculaires qui deviennent alors les nouveaux états propres du système.

$$\psi(r) = C_1\phi(r - R_1) + C_2\phi(r - R_2)$$

En utilisant le formalisme de Dirac et en appelant $1s$ l'état de valence de l'atome d'hydrogène, cette expression s'écrit.

$$|\psi\rangle = C_1|1s\rangle_1 + C_2|2s\rangle_2$$

En portant cette expression de $|\psi\rangle$ dans l'équation de Schrödinger et en multipliant à gauche successivement par $|1s\rangle_1$ et $|1s\rangle_2$ on obtient le système d'équations

$$\begin{aligned}(E_{1s} - E)C_1 - V_2C_2 &= 0 \\ -V_2C_2 + (E_{1s} - E)C_2 &= 0\end{aligned}$$

Où $E_{1s} = {}_1\langle 1s | h | 1s \rangle_1 = {}_2\langle 1s | H | 1s \rangle_2$ est l'énergie de l'électron dans chacun des atomes isolés et $V_2 = -{}_1\langle 1s | H | 1s \rangle_2 = -{}_2\langle 1s | H | 1s \rangle_1$ représente l'énergie de couplage entre les deux atomes. Nous avons négligé l'intégrale de recouvrement : ${}_1\langle 1s | 1s \rangle_2$

Donc la solution du système d'équation est donnée par le déterminant séculaire:

$$\begin{vmatrix} (E_{1s} - E) & -V_2 \\ -V_2 & (E_{1s} - E) \end{vmatrix}$$

C'est-à-dire

$$E_1 = E_{1s} - V_2$$

$$E_a = E_{1s} + V_2$$

En portant ces valeurs dans le système d'équations on obtient les coefficients C_1 et C_2 pour les deux états moléculaires : $C_{1l} = C_{2l} = 1/\sqrt{2}$, $C_{1a} = C_{2a} = 1/\sqrt{2}$, Les orbitales moléculaires sont données en fonction des orbitales atomiques par les expressions:

$$|\psi\rangle_l = \frac{1}{\sqrt{2}} [|1s\rangle_1 + |2s\rangle_2]$$

$$|\psi\rangle_a = \frac{1}{\sqrt{2}} [|1s\rangle_1 - |2s\rangle_2]$$

Les orbitales atomiques $|1s\rangle_1$, $|2s\rangle_2$ et moléculaires $|\psi\rangle_a$, $|\psi\rangle_l$ ainsi que les niveaux d'énergie atomiques et moléculaires sont représentés sur la figure (1.2). L'orbitale $|\psi\rangle_l$ est l'orbitale liante, l'orbitale $|\psi\rangle_a$ est l'orbitale anti-liante

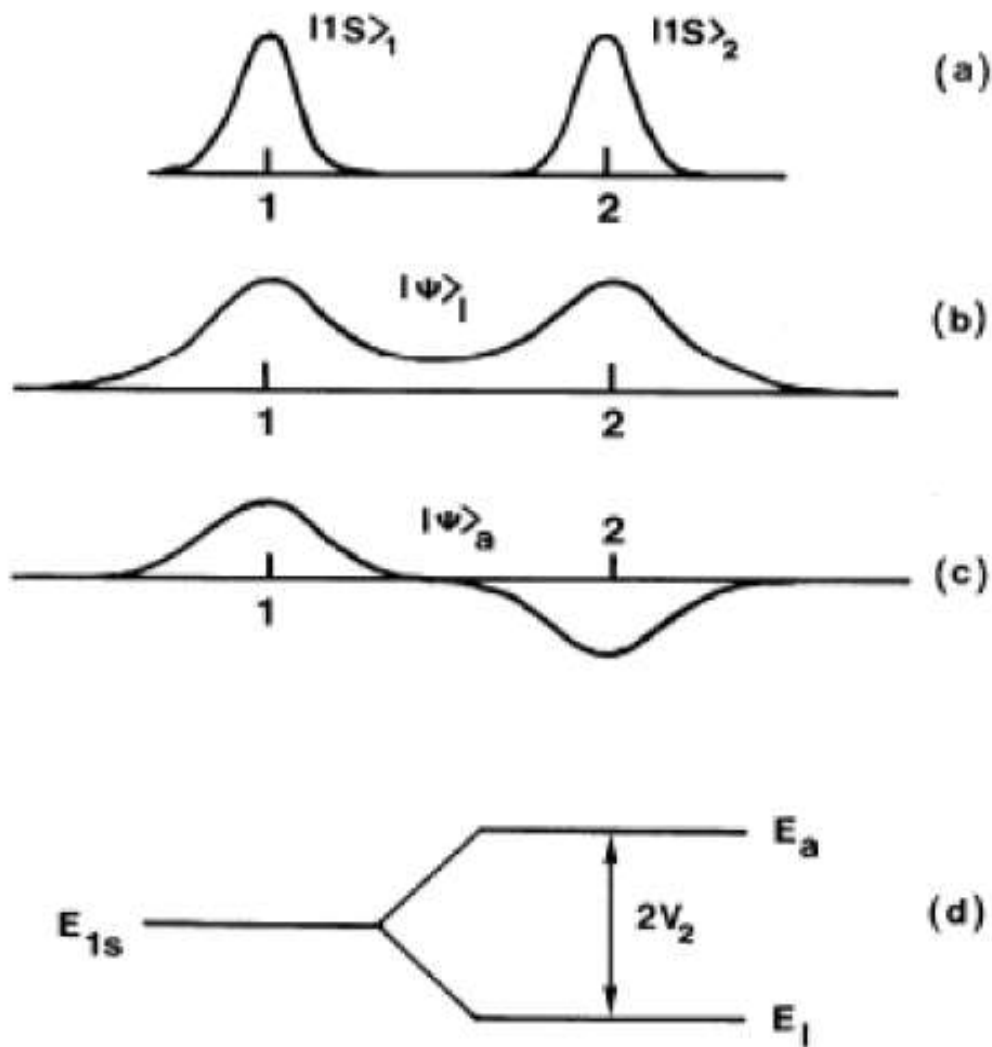


Figure 1.2 : Molécule d'hydrogène, orbitales atomiques et niveaux d'énergie

Il faut noter que dans l'association des deux atomes la densité d'états est conservée et qu'en outre, à l'équilibre thermodynamique, les électrons occupent les états de plus basse énergie. Ainsi les deux états atomiques $|1s\rangle$ dégénérés deux fois compte tenu du spin, donnent naissance à deux états

moléculaires dégénérés deux fois. Les deux électrons occupent l'état liant ψ_{11} , qui est l'état de plus basse énergie.

1.2. Bandes d'énergie

Considérons un atome de silicium Si isolé, les niveaux énergétiques de ses électrons sont discrets (voir le modèle de Bohr pour l'hydrogène). Lorsque l'on rapproche de ce dernier un atome identique, les niveaux énergétiques discrets de ses électrons se scindent en deux sous l'interaction réciproque des deux atomes. Plus généralement, lorsque l'on approche N atomes, les niveaux énergétiques se scindent en N niveaux. Ces N niveaux sont très proches les uns des autres et si la valeur de N est grande, ce qui est le cas pour un cristal, ils forment une bande d'énergie continue. La notion de rapprochement des atomes est donnée par la distance inter-atomique d .

A présent considérons des atomes de silicium Si arrangés aux nœuds d'un réseau périodique, mais avec une maille très grande de telle manière que les atomes puissent être considérés comme isolés. Les deux niveaux les plus énergétiques sont repérés par E_1 et E_2 . Rapprochons homothétiquement les atomes les uns des autres, les états énergétique électronique se scindent et forment deux bandes continues appelées **bande de conduction BC** et **bande de valence BV**. La figure 1.3 montre la formation de ces bandes en fonction de la distance interatomique.

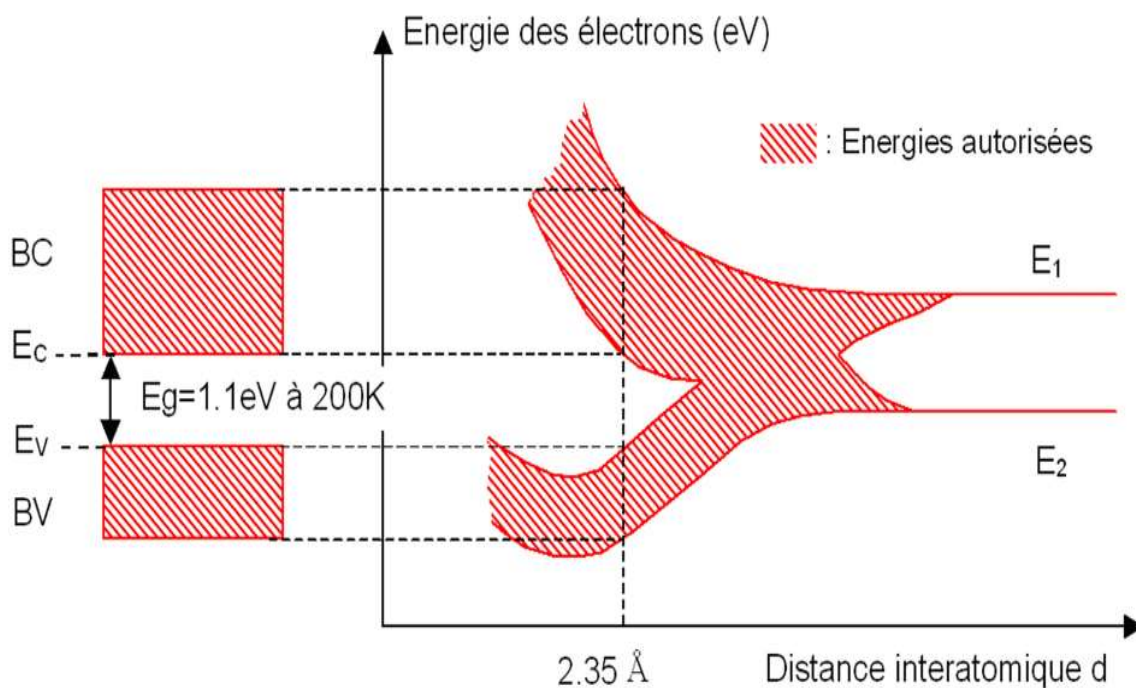


Figure 1.3 : Formation des bandes d'énergie pour les électrons d'atomes de Si arrangés en mailles cristallines de type diamant

Pour les électrons d'un cristal de silicium ($d_0=2,35 \text{ \AA}$), on constate qu'il existe deux bandes continues d'énergie (BC et BV) et que ces bandes sont séparées par une bande interdite car d'énergie inaccessible aux électrons. Cette région interdite est appelée « gap » et sa largeur E_g est caractéristique du matériau. Notons que l'énergie du bas de la bande de conduction est notée E_C et que celle du haut de la bande de valence est notée E_V ainsi nous avons l'égalité $E_g=E_C-E_V$. Précisons que les bandes continues d'énergie BC et BV ne sont qu'une représentation des énergies accessibles par les électrons, ceci ne présage en rien de l'occupation effective de ces bandes par ces derniers.

1.2.1. Cas du solide : Conducteur - Isolant – Semiconducteur

Les matériaux solides peuvent être classés en trois groupes que sont les isolants, les semi-conducteurs et les conducteurs. On considère comme isolants les matériaux de conductivité $\sigma < 10^{-8} \text{ S/cm}$ (diamant 10^{-14} S/cm), comme semi-conducteurs les matériaux tels que $10^{-8} \text{ S/cm} < \sigma < 10^3 \text{ S/cm}$ (par exemple le silicium 10^{-5} S/cm à 10^3 S/cm) et comme conducteurs les matériaux tels que $\sigma > 10^3 \text{ S/cm}$ (Argent: 10^6 S/cm).

Les propriétés électriques d'un matériau sont fonction des populations électroniques des différentes bandes permises. La conduction électrique résulte du déplacement des électrons à l'intérieur de chaque bande. Sous l'action du champ électrique appliqué au matériau l'électron acquiert une énergie cinétique dans le sens opposé au champ électrique.

Considérons à présent une bande d'énergie vide, il est évident de par le fait qu'elle ne contient pas d'électrons, elle ne participe pas à la formation d'un courant électrique. Il en est de même pour une bande pleine. En effet, un électron ne peut se déplacer que si il existe une place libre (un trou) dans sa bande d'énergie. Ainsi, un matériau dont les bandes d'énergie sont vides ou pleines est un **isolant**. Une telle configuration est obtenue pour des énergies de gap supérieures à $\sim 9\text{eV}$, car pour de telles énergies, l'agitation thermique à 300K , ne peut pas faire passer les électrons de la bande de valence à celle de conduction par cassure de liaisons électronique. Les bandes d'énergie sont ainsi toutes vides ou toutes pleines.

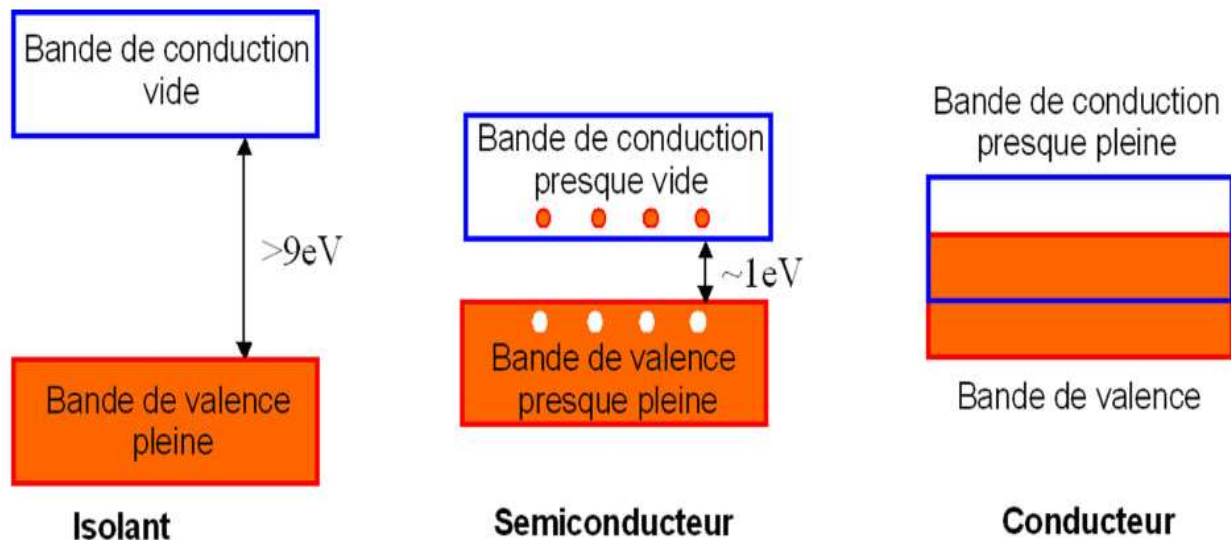


Figure 1.4: Représentation des bandes d'énergie

Un **semi-conducteur** est un isolant pour une température de 0K . Cependant ce type de matériau ayant une énergie de gap plus faible que l'isolant ($\sim 1\text{eV}$), aura de par l'agitation thermique ($T=300\text{K}$), une bande de conduction légèrement peuplée d'électrons et une bande de valence légèrement dépeuplée. Sachant que, la conduction est proportionnelle au nombre d'électrons pour une bande d'énergie presque vide et qu'elle est proportionnelle au nombre de trous pour une bande presque pleine, on déduit que la conduction d'un semiconducteur peut être qualifiée de «mauvaise».

Pour un **conducteur**, l'interpénétration des bandes de valence et de conduction implique qu'il n'existe pas d'énergie de gap. La bande de conduction est alors partiellement pleine (même aux basses températures) et ainsi la conduction du matériau est « élevée ».

1.3. Semi-conducteurs à l'équilibre thermodynamique

Nous allons, dans cet élément, nous concentrer sur les propriétés des semi-conducteurs car ceux-ci ont d'importantes applications. Leur résistivité varie typiquement entre $10^{-3} \Omega.\text{cm}$ et $10^9 \Omega.\text{cm}$, alors que celle des métaux est de l'ordre de $10^{-6} \Omega.\text{cm}$ et celle des isolants peut aller jusqu'à $10^{22} \Omega.\text{cm}$.

1.3.1. Notion de trou

Dans un semi-conducteur, le passage d'un électron de la bande de valence vers la bande de conduction se traduit par l'apparition d'un état inoccupé dans la bande de valence. Ces états inoccupés

ont un rôle très important dans le domaine de la physique des semi-conducteurs ; on les appelle des trous. la conduction électrique d'un semi-conducteur a deux origines :

- les électrons qui se trouvent dans la bande de conduction,
- les « vides » ou trous présents dans la bande de valence. En effet, le mouvement des électrons restants se traduit très simplement par un mouvement d'électrons manquants, i.e. de trous. Ceci est illustré dans la figure 1.5.

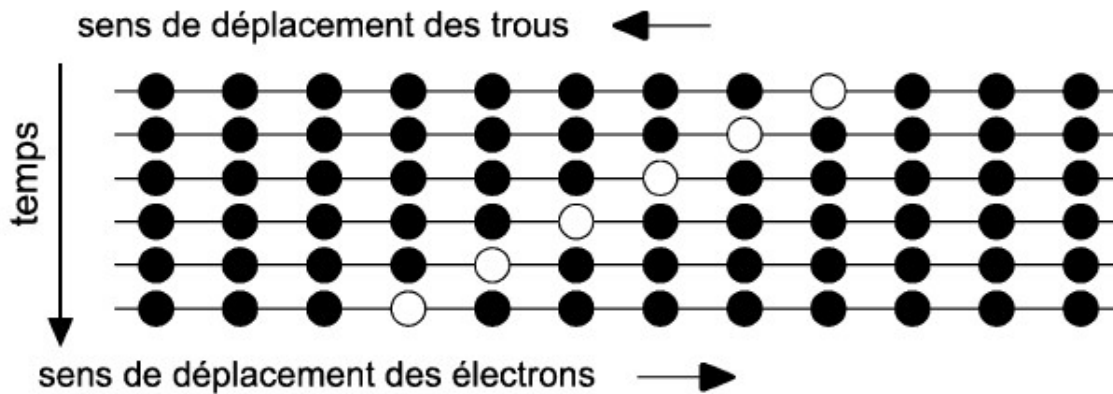


Figure 1.5. Propagation, sur un réseau unidimensionnel composé d'électrons (cercles pleins), d'un trou (cercle vide). Si chaque électron se déplace de gauche à droite d'une position, le trou, qui est une place vacante, se déplace de droite à gauche.

1.3.2. Notion de gap

La bande de valence d'un semi-conducteur est, à température nulle, la dernière bande occupée par les électrons. La bande de conduction est la première bande vide à température nulle. Le gap est l'énergie minimum séparant la bande de valence de la bande de conduction. Il dépend de la température et varie d'environ 10 % entre 0 K et la température ordinaire. Cette dépendance a deux origines essentielles :

- La dilatation thermique du réseau cristallin qui modifie le potentiel périodique, donc la structure de bandes et la valeur du gap en énergie ;
- Les vibrations du réseau car elles dépendent de la température et modifient aussi la structure de bandes, donc la valeur du gap.

Ces deux effets sont de grandeur comparable et conduisent à une variation linéaire en T sauf à basse température où cette variation est quadratique. Ainsi, pour le silicium, on a $E_g(T = 0) = 1,17 \text{ eV}$ et $E_g(T = 300 \text{ K}) = 1,12 \text{ eV}$, avec une variation linéaire au-dessus de 200 K.

1.3.2.1. Nature du Gap

Pour un cristal semi-conducteur tridimensionnel, le maximum de la bande de valence et le minimum de la bande de conduction sont caractérisés par une énergie E et un vecteur d'onde k . Si, dans l'espace réciproque, ce maximum et ce minimum correspondent à la même valeur de $\hbar k$ (cas de l'AsGa), on dit que le semi-conducteur est à gap direct. Pour d'autres cristaux, le silicium en particulier, ce maximum et ce minimum correspondent à des valeurs de k différentes et le semi-conducteur est dit à gap indirect. La nature du gap est importante pour les application opto-électroniques qui mettent en jeu à la fois des électrons et des photons. En effet, lors de la transition d'un électron de la bande de valence vers la bande de conduction — donc avec création d'un trou dans la bande de valence et d'un électron dans la bande de conduction — ou lors de la recombinaison d'un électron de la bande de conduction avec un trou de la bande de valence, il faut conserver :

- l'énergie (relation scalaire),
- et l'impulsion (relation vectorielle).

Lorsqu'un semi-conducteur est à gap direct, la transition d'un électron de la bande de valence vers la bande de conduction se fait à quantité de mouvement $\hbar k$ constante puisque l'impulsion d'un photon est négligeable devant celle de l'électron. La conservation de l'impulsion est donc automatiquement assurée. Cette transition s'observe, par exemple, lorsque le semi-conducteur absorbe un photon de longueur d'onde :

$$\lambda(\text{nm}) = \frac{1,240}{E_g} (\text{eV})$$

Où E_g est la valeur du gap. De même, un électron de la bande de conduction peut facilement se recombiner avec un trou en émettant un photon d'énergie λ ; c'est ce que l'on appelle une transition radiative. Ce phénomène est utilisé dans les diodes électroluminescentes et les diodes laser. Le mécanisme de transition directe est schématisé dans la partie gauche de la figure 1.6.

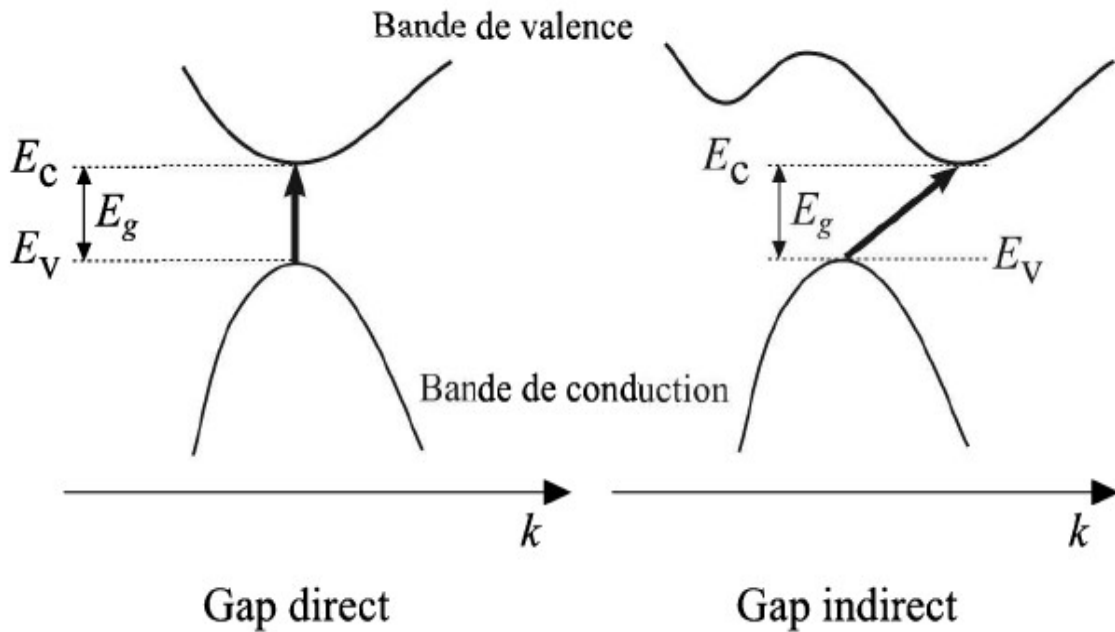


Figure 1.6. Différence entre une transition directe et une transition indirecte.

Pour un semi-conducteur à gap indirect, comme le silicium, les choses sont plus complexes.

Pour qu'un électron de conduction, d'impulsion $\hbar k$, se recombine avec un trou d'impulsion nulle de la bande de valence, il faut, pour conserver l'impulsion, faire intervenir les phonons du réseau (l'énergie du phonon créé est typiquement de l'ordre 0,01 à 0,03 eV). Un phonon participe au processus de recombinaison pour assurer la conservation de l'impulsion (grâce à l'impulsion du réseau cristallin). Ceci rend la combinaison radiative beaucoup moins probable que dans le cas des semi-conducteurs à gap direct. Par conséquent, le silicium, tout comme le germanium qui est aussi un semi-conducteur à gap indirect, est un très mauvais émetteur de lumière. Tout ceci s'applique aussi à l'excitation d'un électron de la bande de valence vers la bande de conduction. Ainsi, dans le cas d'un gap indirect, on ne peut exciter optiquement le semi-conducteur, à $T = 0$ K, que s'il intervient un phonon de vecteur d'onde $\mathbf{K}$.

Le mécanisme de transition indirecte a donc une plus faible probabilité de se produire que le processus direct, car il met en jeu 3 particules (l'électron, le photon et le phonon) au lieu de deux (l'électron et le photon). Ceci explique en particulier pourquoi l'AsGa est un bon matériau pour les applications opto-électroniques (i.e. mettant en jeu des photons et des électrons) alors que le silicium est moins bon. Lorsque la température augmente, depuis 0 K, le nombre de phonons augmente fortement et la probabilité d'excitation optique croît aussi.

1.3.2.2. Masse effective et relation de dispersion

Pour un semi-conducteur à l'équilibre ou voisin de celui-ci, les niveaux inoccupés sont proches du sommet de la bande de valence. De même, les niveaux occupés dans la bande de conduction sont voisins du bas de celle-ci. Ceci permet de faire une approximation parabolique (harmonique) de ces bandes au voisinage de leur extremum.

Considérons la bande de conduction. Soit $\mathbf{k}_0$ la position du minimum de $E(k)$. L'approximation harmonique consiste à remplacer la bande de valence au voisinage de $\mathbf{k}_0$ par le paraboloïde oscillateur en $\mathbf{k}_0$. Supposons, pour simplifier la discussion, le système suffisamment symétrique pour que ce paraboloïde soit de révolution. L'approximation harmonique est un développement de Taylor jusqu'au second ordre en k , soit :

$$E(\mathbf{k}) \approx E(\mathbf{k}_0) + \left(\frac{d^2 E(\mathbf{k})}{d\mathbf{k}^2} \right)_{\mathbf{k}=\mathbf{k}_0} (\mathbf{k} - \mathbf{k}_0)^2 = E(\mathbf{k}_0) + A(\mathbf{k} - \mathbf{k}_0)^2$$

Où A est une constante. Le fait que $E(k)$ soit minimum en $k = k_0$ entraîne que $A > 0$ et que le terme linéaire du développement soit nul. On peut alors définir la masse effective de l'électron par la relation :

$$\frac{\hbar^2}{2m_e^*} = A = \left(\frac{d^2 E(\mathbf{k})}{d\mathbf{k}^2} \right)_{\mathbf{k}=\mathbf{k}_0}$$

Soit :

$$\left(\frac{1}{m_e^*} \right) = \frac{2}{\hbar^2} \frac{d^2 E(\mathbf{k})}{d\mathbf{k}^2}$$

On peut écrire

$$E(\mathbf{k}) \approx E(\mathbf{k}_0) + \frac{\hbar^2}{2m_e^*} (\mathbf{k} - \mathbf{k}_0)^2$$

La vitesse de groupe au voisinage de $\mathbf{k} = \mathbf{k}_0$ vaut alors :

$$\mathbf{v}_g = \frac{1}{\hbar} \frac{d}{d\mathbf{k}} E(\mathbf{k}) \approx \frac{\hbar (\mathbf{k} - \mathbf{k}_0)}{m_e^*}$$

Lorsque $k_0 = 0$, ce qui est le cas de l'As Ga par exemple, on retrouve les expressions de l'électron libre à ceci près que la masse de celui-ci, m , est remplacée par la masse effective m_e^* . La masse effective est donc associée à la courbure de $E(k)$ au voisinage du minimum. Si cette courbure, appelée aussi rigidité dans le cas d'un oscillateur harmonique, est grande, la masse effective est petite et réciproquement. Ceci est illustré sur la figure 1.7.

Nous avons considéré une bande de conduction suffisamment symétrique pour pouvoir utiliser la relation

$$E(\mathbf{k}) \approx E(\mathbf{k}_0) + \frac{\hbar^2}{2m_e^*} (\mathbf{k} - \mathbf{k}_0)^2$$

Ce n'est souvent pas le cas et toutes les directions de l'espace des $\mathbf{k}$ ne sont pas équivalentes en partant de $\mathbf{k}_0$. Dans cette situation, le développement de Taylor jusqu'au second ordre s'écrit :

$$E(\mathbf{k}) \approx E(\mathbf{k}_0) + \sum_{i,j} \left(\frac{d^2 E(\mathbf{k})}{dk_i dk_j} \right)_{\mathbf{k}=\mathbf{k}_0} (k_i - k_{0i})(k_j - k_{0j})$$

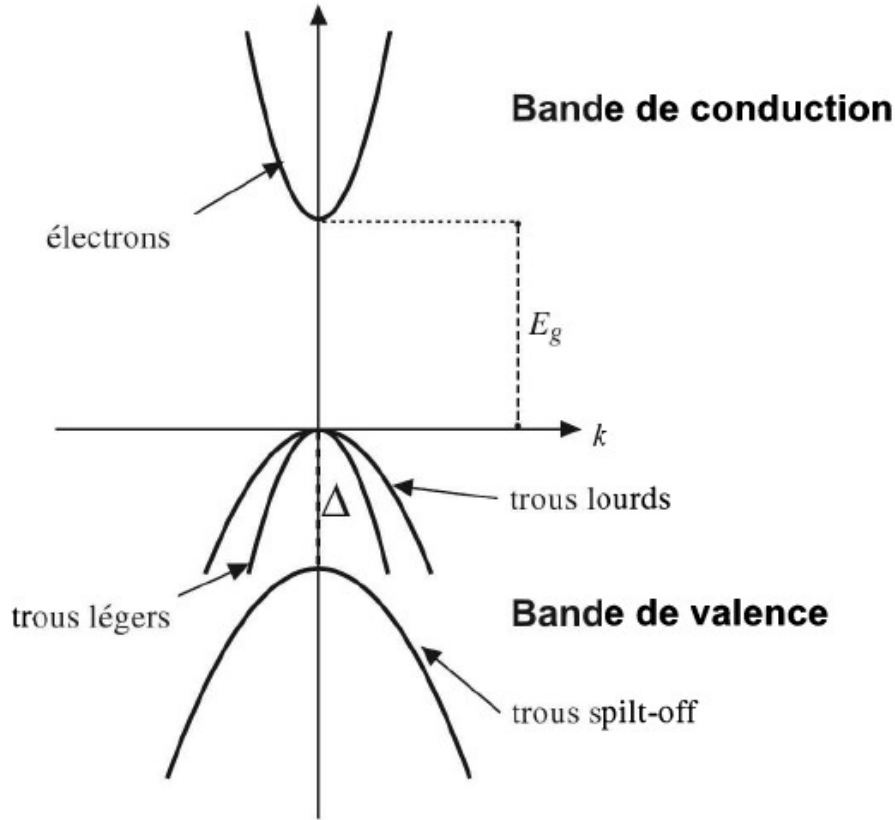


Figure 1.7: Illustration très simplifiée des bords de bandes pour un semi-conducteur à gap direct.

Pour la bande de valence, il y a deux bandes dégénérées qui correspondent aux trous légers et lourds. Une troisième bande, située plus bas en énergie, correspond à une levée de dégénérescence due à l'interaction spin-orbite; ce sont les trous split-off. Le couplage spin-orbite peut parfois conduire à un écart, Δ , entre les sous-niveaux, supérieur au gap (dans InSb, par exemple, $\Delta = 0,9$ eV alors que $E_g = 0,23$ eV). Ce n'est pas le cas pour le silicium pour lequel $\Delta = 0,04$ eV alors que $E_g = 1,12$ eV à la température ordinaire.

Où les indices i, j repèrent les composantes x, y et z des vecteurs d'onde. La masse effective est définie par la relation :

$$\left(\frac{1}{m_e^*} \right)_{ij} = \frac{2}{\hbar^2} \left(\frac{d^2 E(\mathbf{k})}{dk_i dk_j} \right)$$

1.4. Structure des bandes d'énergie

La structure réelle des bandes d'énergie est beaucoup plus complexe que les simples schémas que l'on a l'habitude de présenter pour comprendre les propriétés des semiconducteurs. Néanmoins, ceux-ci sont suffisants pour comprendre la physique de base mise en jeu. Dans cette section, nous allons présenter la structure des bandes d'énergie de deux semi-conducteurs, celle de l'arséniure de

gallium, qui est une bonne illustration d'un semi-conducteur à gap direct, et celle du silicium, qui est un semiconducteur à gap indirect.

1.4.1. Arséniure de gallium

L'AsGa est un semi-conducteur à gap direct. Le minimum de la bande de conduction est situé en $k(0, 0, 0)$, i.e. au point Γ . Au voisinage de ce point, les surfaces d'énergie constante sont des sphères centrées en Γ . Pour cette raison, on qualifie l'AsGa de semi-conducteur univallée. Ce schéma est observé pour la plupart des semi-conducteurs III—V et II—VI. La structure de bandes de l'AsGa est indiquée sur la figure 1.8. On retrouve, trois bandes correspondant respectivement aux trous légers ($m_h^* \approx 0,1m$), aux trous lourds ($m_h^* \approx 0,45m$) et à la bande « split-off » ($m_h^* \approx 0,2m$). La valeur assez grande du déplacement spin orbite ($\Delta \approx 0,34$ eV) fait que cette dernière n'intervient en général pas dans les propriétés électroniques et optoélectroniques de l'AsGa.

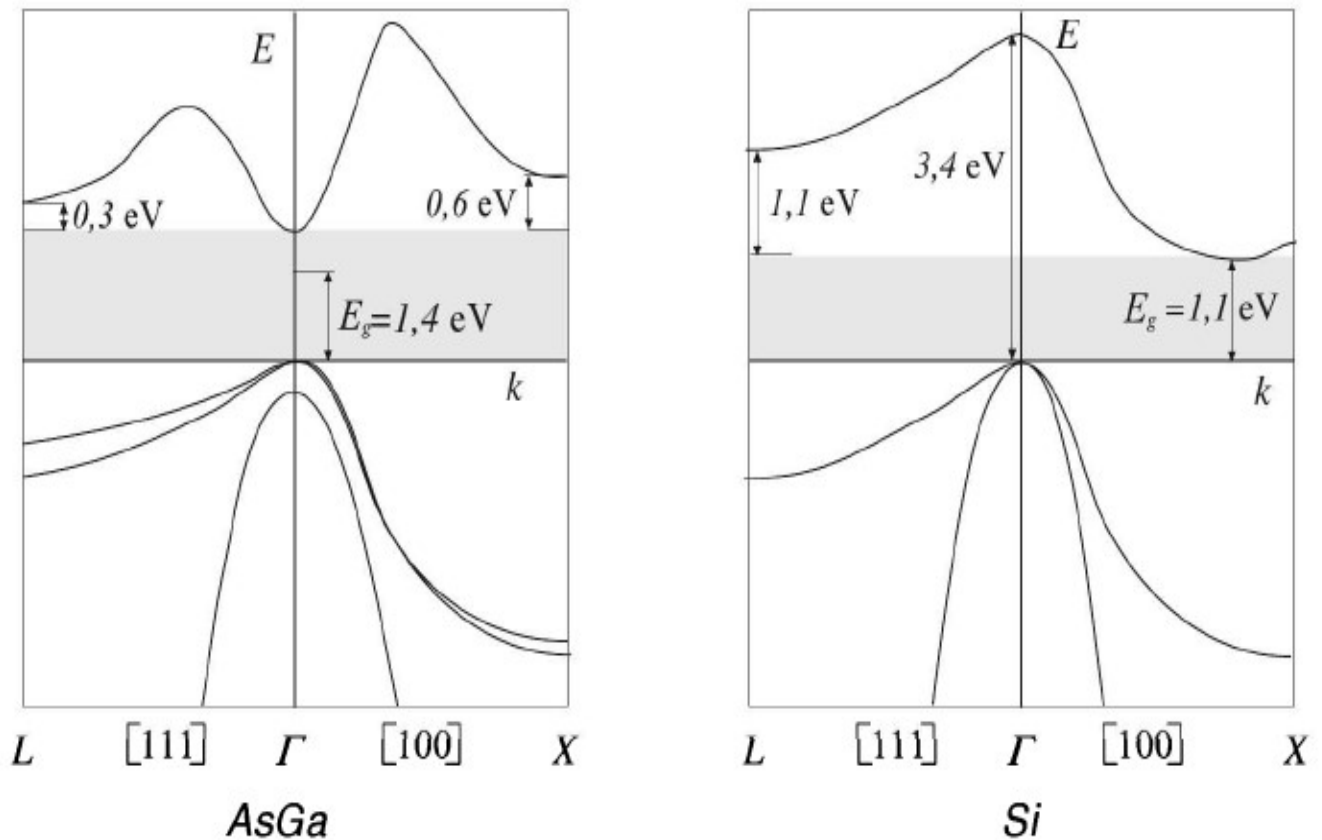


Figure 1.8: Structure simplifiée des bandes d'énergie de l'AsGa selon les directions $[100]$ et $[111]$ du cristal. Le point Γ est situé en $k = (0,0,0)$, le point X en $k = 2\pi/a(1,0,0)$ et le point L en

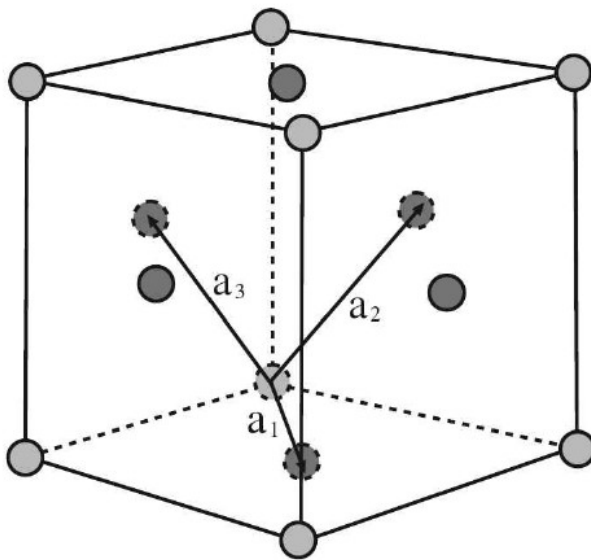
$k = \pi/a(1,1,1)$. Le gap de l'AsGa est direct et celui du silicium indirect

Dans le diagramme de la figure 1.8, on remarque l'existence de deux minima dont l'énergie est supérieure à celle du point Γ . Il s'agit de la vallée située en L et de celle située en X ; elles sont

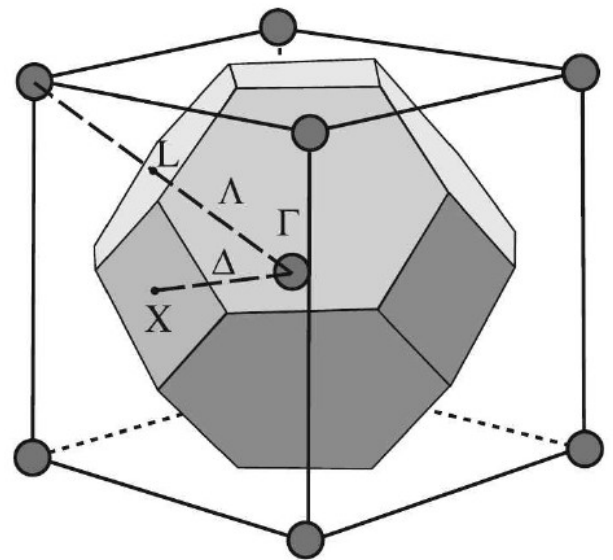
respectivement positionnées à 0,3 eV et 0,58 eV au-dessus de la vallée Γ . Par suite de la symétrie du cristal, il y a 4 vallées équivalentes de type L .

1.4.2. Silicium

Le minimum de la bande de conduction est situé au point d'abscisse $(0,0,1k_0 = 0,85 k_x)$, où k_x est l'abscisse du point X de la direction Δ à de la figure 1.8 (direction $[100]$). Par suite de la symétrie du cristal, il existe 6 directions équivalentes qui sont $[100]$, $[-100]$, $[010]$, $[0-10]$, $[001]$ et $[00-1]$. Le silicium présente donc des minima égaux en énergie dans ces directions. On qualifie pour cette raison le silicium de semi-conducteur multivallées. Au voisinage des minima, la variation de $E(\mathbf{k})$ n'est pas isotrope ; elle est plus rapide dans le plan perpendiculaire à l'axe Δ . Les surfaces d'énergie constante au voisinage des minima sont des paraboloides de révolution autour des axes Δ .



réseau direct (cfc)



réseau réciproque (cc)

Figure 1.9. Structure cfc à gauche (réseau direct); structure cubique centrée à droite (réseau réciproque). Quelques points et lignes de symétrie citées dans le texte sont indiqués et la première zone de Brillouin est le polyèdre en grisé.

1.5. Densité d'états

La densité d'état intervient dans le calcul de nombreuses quantités physiques des semiconducteurs. Elle est définie comme le nombre de **micro-états par unité de volume et d'énergie** ; nous la noterons

$\rho(E)$. Nous allons calculer $\rho(E)$ pour un gaz d'électrons indépendants. Il nous sera ensuite aisé d'utiliser les expressions obtenues pour les porteurs de charge, dont l'énergie est proche du bord des bandes, en utilisant la notion de masse effective. Pour des électrons libres :

$$E = \frac{\hbar^2 k^2}{2m} \implies dE = \frac{2\hbar^2 k dk}{2m} = \hbar \sqrt{\frac{2E}{m}} dk$$

Le nombre de micro-états dans le volume V ayant leur vecteur d'onde compris entre $\mathbf{k}$ et $\mathbf{k}+d\mathbf{k}$ est obtenu en évaluant le volume dans l'espace de phase et en divisant par h^3 . Ce volume est obtenu en considérant l'espace compris entre deux sphères de rayon p et $p + dp$ (rappelons que $p = \hbar k$) que l'on multiplie par le volume V de l'espace ordinaire. Il faut ensuite tenir compte, dans le dénombrement, de la dégénérescence du spin (= 2). On obtient :

$$2 \times \frac{4\pi p^2 dp}{h^3} V = 2 \times \frac{4\pi k^2 dk}{(2\pi)^3} V = \frac{k^2 dk}{\pi^2} V$$

Le nombre de micro-états compris entre E et $E + dE$ et la densité d'état valent :

$$\rho(E)dE = \frac{k^2 dk}{\pi^2} = \frac{\sqrt{2}m^{3/2}}{\pi^2 \hbar^3} \sqrt{E} dE$$

Où

$$\rho(E) = \frac{\sqrt{2}m^{3/2}}{\pi^2 \hbar^3} \sqrt{E}$$

Les électrons sont parfois confinés dans des structures à 2 dimensions ou à une dimension. Il est intéressant d'évaluer $\rho(E)$ dans ces deux cas. À deux dimensions, on a :

$$\begin{aligned} \rho(E)dE &= 2 \times \frac{2\pi k dk}{(2\pi)^2} = 2 \times \frac{2m}{(2\pi)^2 \hbar^2} dE \\ \implies \rho(E) &= \frac{m}{\pi \hbar^2} \end{aligned}$$

À une dimension :

$$\rho(E)dE = 2 \times \frac{dk}{2\pi} = 2 \times \sqrt{\frac{m}{2E}} \frac{1}{\pi \hbar} dE$$

$$\Rightarrow \rho(E) = \frac{\sqrt{2m}}{\pi\hbar} \frac{1}{\sqrt{E}}$$

La dépendance de la densité d'états en fonction de l'énergie est différente selon la dimensionnalité du système. À trois dimensions, elle varie comme $E^{1/2}$, à deux dimensions elle est constante et à une dimension elle varie comme $E^{-1/2}$.

1.5.1. Densité d'états d'électrons et trous

Nous allons maintenant considérer un semi-conducteur à la température T . Soit E_c l'énergie minimum de la bande de conduction et E_v l'énergie maximum de la bande de valence. À $T = 0$ K, tous les électrons sont dans la bande de valence. Lorsque T augmente, certains d'entre-eux, ceux qui ont une énergie proche de E_v , passent dans la bande de conduction, au voisinage de E_c . En effet, les excitations les plus probables sont celles qui demandent le moins d'énergie, donc celles qui sont à peine supérieures à $E_c - E_v$. Pour calculer n_e , le nombre d'électrons par unité de volume dans la bande de conduction :

- On évalue le nombre $\rho(E)dE$ de micro-états dont le vecteur d'onde est compris entre E et $E + dE$.
 - On multiplie $\rho(E)dE$ par la probabilité $f(E)$ d'occupation d'un état par un électron.
 - Enfin, on intègre sur l'espace E correspondant à la bande de conduction (BC), i.e. sur $E > E_c$.
- Ceci conduit à :

$$n_e = \int_{BC} \rho(E) f(E) dE = \int_{E_c}^{\infty} \rho(E) f(E) dE$$

Le raisonnement est analogue pour calculer n_h , le nombre de trous par unité de volume dans la bande de valence, à ceci près que la probabilité d'obtenir un trou d'énergie E est $1 - f(E)$ et que l'intégration se fait sur la bande de valence (BV) :

$$n_h = \int_{BV} \rho(E) (1 - f(E)) dE$$

La probabilité d'occupation d'un niveau, $f(E)$, est donnée par la distribution de Fermi-Dirac :

$$f(E) = \frac{1}{1 + \exp[(E - E_F)/k_B T]}$$

Rappelons encore une fois que E_F , dénommée énergie de Fermi dans le domaine des semi-conducteurs, est en fait le potentiel chimique du système.

La densité d'état dans la bande de conduction s'obtient à partir de la forme de $\rho(E)$ obtenue dans la section 1.4.4. (équation):

$$\rho(E)dE = \frac{k^2 dk}{\pi^2} = \frac{\sqrt{2}m^{3/2}}{\pi^2 \hbar^3} \sqrt{E} dE \quad \text{ou} \quad \rho(E) = \frac{\sqrt{2}m^{3/2}}{\pi^2 \hbar^3} \sqrt{E}$$

En utilisant la masse effective de l'électron pour tenir compte des effets de milieu. Celle-ci est supposée constante puisque l'énergie des électrons est proche de E_c . Par ailleurs, dans l'équation précédente, il faut remplacer E par $E - E_c$, puisque l'origine des énergies de la bande de conduction est en E_c . L'expression de n_e , est donc :

$$n_e = \frac{\sqrt{2} (m_e^*)^{3/2}}{\pi^2 \hbar^3} \int_{E_c}^{\infty} \frac{(E - E_c)^{1/2} dE}{\exp\left(\frac{E - E_F}{k_B T}\right) + 1}$$

Pour un semi-conducteur intrinsèque (i.e. pur), l'énergie de Fermi est loin du bord de la bande de conduction, entre E_v et E_c . Ceci implique que $\exp((E - E_F)/k_B T) \gg 1$. On peut donc négliger l'unité devant l'exponentielle ; c'est l'approximation de Boltzmann basée sur le fait qu'il y a peu d'électrons (de trous) dans la bande de conduction (de valence) par rapport au nombre de micro-états disponibles. Elle est valable lorsque n_e , est petit, typiquement inférieur à 10^{16} cm^{-3} . Compte tenu de ces hypothèses, on déduit :

$$n_e = 2 \left(\frac{m_e^* k_B T}{2\pi \hbar^2} \right)^{3/2} \exp\left(\frac{E_F - E_c}{k_B T}\right) = N_e \exp\left(\frac{E_F - E_c}{k_B T}\right)$$

Où l'on a introduit la quantité

$$N_e = 2 \left(m_e^* k_B T / 2\pi \hbar^2 \right)^{3/2}$$

qui est la densité effective d'états au bord de la bande de conduction. C'est le nombre par lequel il faut multiplier la probabilité d'occupation du niveau pour obtenir n_e . En fait, les hypothèses que nous avons faites reviennent à considérer les électrons de la bande de conduction comme un gaz parfait de particules de masse m_e^* , d'énergie E_c , à la température T . Le facteur **2**, au début de la deuxième ligne de l'équation précédent, provient de la dégénérescence liée au spin des électrons.

On obtient, par un raisonnement tout à fait analogue, le nombre de trous par unité de volume dans la bande de valence en partant de l'équation:

$$n_h = \int_{BV} \rho(E) (1 - f(E)) dE$$

et en faisant les mêmes hypothèses que ci-dessus qui sont valables pour les semi-conducteurs intrinsèques.

Cela donne :

$$n_h = 2 \left(\frac{m_h^* k_B T}{2\pi \hbar^2} \right)^{3/2} \exp \left(-\frac{E_F - E_V}{k_B T} \right) = N_h \exp \left(-\frac{E_F - E_V}{k_B T} \right)$$

Où

$$N_h = 2 \left(m_h^* k_B T / 2\pi \hbar^2 \right)^{3/2}$$

est la densité effective d'états au bord de la bande de valence. Pour un porteur de charge de masse effective m^* (m_e^* ou m_h^* , nous avons, si m est la masse de l'électron libre et T la température évaluée en kelvins, l'expression suivante pour N (N_e ou N_h) :

$$N = 2,5 \left(\frac{m^*}{m} \right)^{3/2} \left(\frac{T}{300 \text{ K}} \right)^{3/2} \times 10^{19} \text{ cm}^{-3}$$

Comme les facteurs exponentiels intervenant dans le calcul du nombre de porteurs de charge sont toujours inférieurs à l'unité, le nombre maximum de porteurs, dans le cas non dégénéré, est toujours inférieur à 10^{18} - 10^{19} cm^{-3} . Dans le calcul du nombre de porteurs de charge, il faut prendre les masses effectives calculées à partir des densités d'états.

Chapitre 2.

Semi-conducteurs Intrinsèques et Extrinsèques

2.1. Un semi-conducteur est dit « intrinsèque » ?

Un semi-conducteur est dit « intrinsèque » s'il est à l'état pur. Dans ce cas, le nombre d'électrons dans la bande de conduction est **égal** au nombre de trous dans la bande de valence. Par conséquent, $n_e = n_h$ et l'on déduit à partir des équations au-dessous :

$$n_e = 2 \left(\frac{m_e^* k_B T}{2\pi \hbar^2} \right)^{3/2} \exp \left(\frac{E_F - E_C}{k_B T} \right) = N_e \exp \left(\frac{E_F - E_C}{k_B T} \right)$$

$$n_h = 2 \left(\frac{m_h^* k_B T}{2\pi \hbar^2} \right)^{3/2} \exp \left(-\frac{E_F - E_V}{k_B T} \right) = N_h \exp \left(-\frac{E_F - E_V}{k_B T} \right)$$

L'énergie de Fermi du système intrinsèque, E_F^i vaut :

$$E_F = \frac{E_C + E_V}{2} + \frac{3}{4} k_B T \text{Log} \left(\frac{m_h^*}{m_e^*} \right)$$

Le niveau de Fermi d'un système intrinsèque est donc proche du milieu de la bande interdite. Dans l'hypothèse, souvent faite, selon laquelle la masse effective des électrons est égale à celle des trous ($m_e^* = m_h^*$), on a :

$$E_F \simeq (E_C + E_V) / 2$$

Cela indique que l'énergie de Fermi est située au milieu de la bande interdite. Pour de nombreux semi-conducteurs (mais pas le silicium), $m_h^* > m_e^*$. Leur niveau de Fermi est situé légèrement **au-dessus** du milieu de la bande interdite. Posons $n_i = n_e = n_h$. On a, en utilisant les deux équations correspondant au nombre des électrons et les trous dans la BC et la BV respectivement.

$$\begin{aligned} n_i^2 &= n_e n_h = N_e N_h \exp(-E_g / k_B T) \\ &= 4 \left(\frac{k_B T}{2\pi \hbar^2} \right)^3 (m_e^* m_h^*)^{3/2} \exp(-E_g / k_B T) \end{aligned}$$

La concentration de porteurs intrinsèques vaut par conséquent:

$$n_i = 2 \left(\frac{k_B T}{2\pi \hbar^2} \right)^{3/2} (m_e^* m_h^*)^{3/4} \exp(-E_g/2k_B T)$$

Cette relation montre que la concentration intrinsèque dépend du gap E_g et qu'elle augmente avec la température comme $T^{3/2}$. On retrouve le fait que la conductivité d'un semi-conducteur augmente avec la température. L'équation précédent conduit à :

$$n_i = 2,5 (m_e^*/m)^{3/4} (m_h^*/m)^{3/4} (T/300 \text{ K})^{3/2} \exp(-E_g/2k_B T) \times 10^{19} \text{ cm}^{-3}$$

Le nombre de porteurs est faible à température ordinaire (10^{10} par cm^{-3} pour le silicium) comparé au nombre d'électrons dans les métaux (typiquement de l'ordre de 10^{21} par cm^{-3}).

$$n_i^2 = n_e n_h$$

C'est la loi d'action de masse. Elle est satisfaite quelles que soient les valeurs de n_e et n_h .

2.2. Dopage

Les propriétés électriques d'un cristal semi-conducteur sont profondément modifiées si l'on remplace certains atomes du réseau par des atomes ayant, par rapport à l'atome substitué, un électron en plus ou en moins dans son cortège électronique. On désigne ceci sous le nom de dopage. Ce dernier consiste à introduire des impuretés dans le cristal qui vont avoir pour conséquence de créer des niveaux d'énergie dans la bande interdite. Ils laissent aussi une densité charge qui se traduit par un champ électrique et parfois une barrière de potentiel à l'intérieur du semiconducteur. De telles propriétés sont absentes dans un métal ou un isolant.. Un semi-conducteur dopé est appelé semi-conducteur extrinsèque.

Dans tous les cas, la concentration du dopant est faible, 10^{16} ou 10^{17} cm^{-3} pour donner un exemple. De plus, chaque dopant a une concentration limite qui dépend de sa nature, du cristal semi-conducteur à doper et de la température. Par exemple, cette concentration limite est de l'ordre de 10^{20} cm^{-3} pour du bore (B) dans du silicium à 1000°C . Le P et l'As ont des limites à peine supérieures dans les mêmes conditions. Pour des concentrations plus grandes, il se forme des précipités qui rendent le semiconducteur impropre à une utilisation en électronique. Plusieurs méthodes sont possibles pour réaliser ce dopage :

On peut doper le silicium fondu avant la croissance du cristal.

- *l'implantation ionique permet de bien contrôler le dopage. La méthode consiste à bombarder le cristal avec des ions dont l'énergie cinétique est de l'ordre de 15 à 200 keV (éventuellement jusqu'au MeV).*
- *Les implanteurs d'ions sont de petits accélérateurs de particules qui permettent de communiquer de l'énergie cinétique aux ions de l'atome dopant. Lors du dopage du cristal, on prend soin d'éliminer*

l'effet de canalisation en orientant la plaquette de silicium d'un angle de 7° par rapport à la direction incidente du faisceau d'ions. L'étape d'implantation est suivie d'un recuit thermique pour restaurer le réseau cristallin qui s'est amorphisé localement aux points d'implantation.

- *On peut faire passer sur la surface du silicium un gaz contenant une impureté dopante comme BBr_3 ou $POCl_3$, par exemple. La diffusion du dopant se fait alors dans un four à haute température 800-1000°C).*
- *D'autres méthodes particulières peuvent être utilisées comme le bombardement par des neutrons qui permettent de doper du Si avec du phosphore selon une réaction nucléaire de transmutation.*

L'évolution vers une réduction de taille des composants des circuits intégrés conduit à réaliser des dopages de moins en moins profonds ce qui permet l'utilisation de techniques laser pour le recuit ou le dopage en phase gazeuse assistée par laser.

Prenons le cas du silicium situé dans la colonne IV de la classification périodique. On peut introduire dans son cristal des impuretés (appelées dopants) appartenant à la colonne III (du bore, par exemple) ou V (du phosphore, par exemple). Ces dopants doivent se placer en position substitutionnelle dans le cristal (ils doivent prendre la place d'un atome de silicium). En effet, les dopants placés en position interstitielle (dans un interstice du réseau) ne conduisent pas à une modification notable des propriétés électriques.

Un semi-conducteur intrinsèque possède des impuretés résiduelles qui peuvent jouer le rôle d'accepteurs ou de donneurs d'électrons. On peut le doper avec des impuretés de type opposés pour compenser l'effet des impuretés résiduelles. Dans ce cas on dit que le semi-conducteur est **compensé**.

2.2.1. Exemple de dopage

La figure (2.1) montre trois représentations de liaison de base d'un semiconducteur. La figure (2.1.a) montre du silicium intrinsèque, qui est très pur et contient une quantité négligeable de impuretés. Chaque atome de silicium partage ses quatre électrons de valence avec les quatre voisins formant quatre liaisons covalentes La figure (2.1.b) montre du silicium de type ***n***, où un atome de phosphore en substitution avec cinq électrons de valence a remplacé un atome de silicium et un électron chargé négativement est donné au réseau dans la bande de conduction. L'atome de phosphore est appelé un donneur (d'électrons).

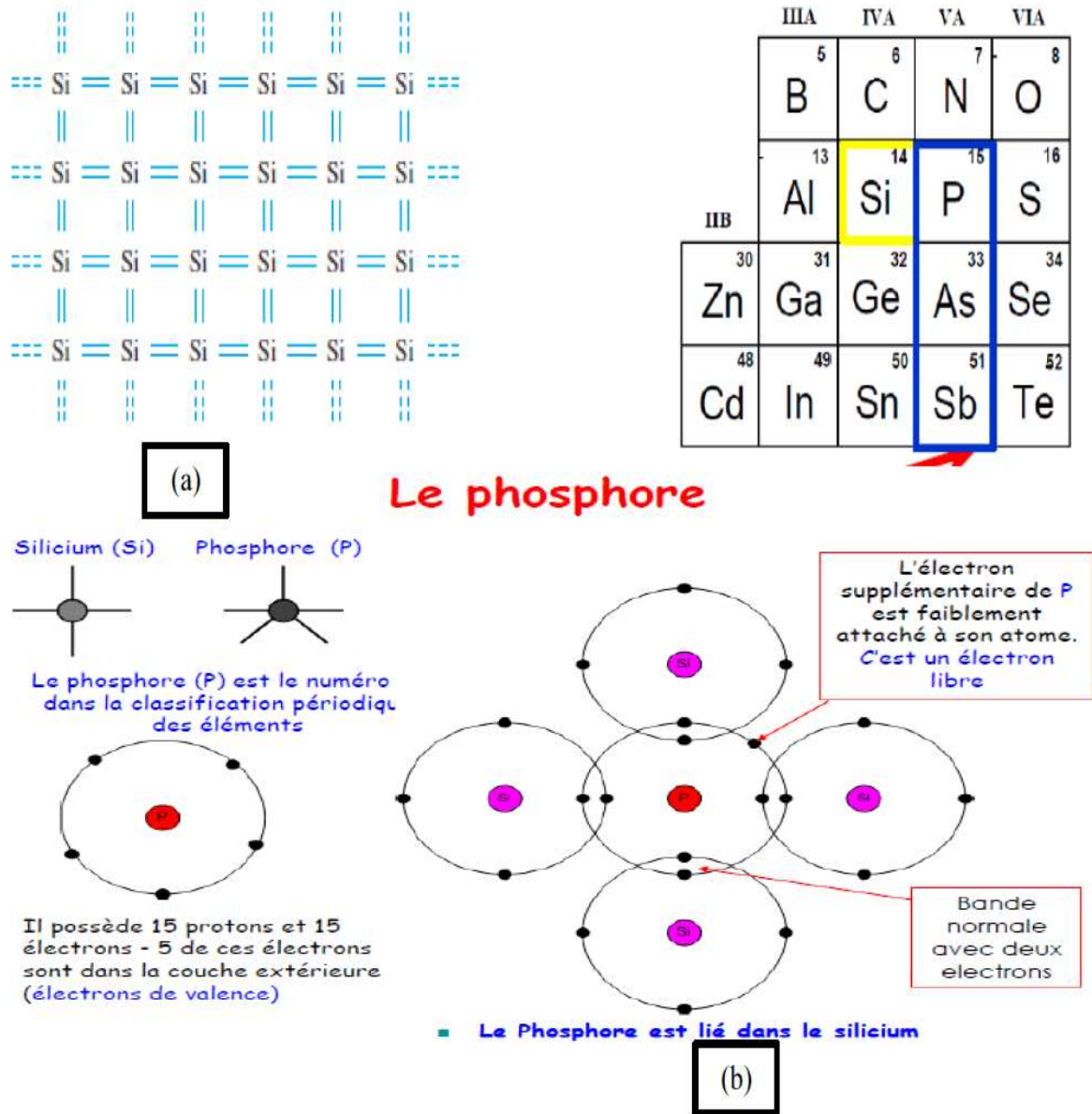
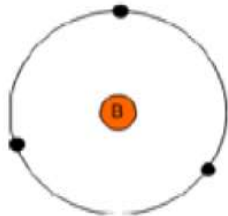


Figure 2.1 : un atome silicium, un atome de phosphore (donneur)

La figure (2.2) de même montre que lorsqu'un atome de bore avec trois électrons de valence remplace un atome de silicium, un trou chargé positivement est créé dans la bande de valence et un électron supplémentaire sera accepté pour former quatre liaisons covalentes autour du bore. On obtient un semiconducteur de type *p*, et le bore est un atome accepteur (d'électrons).

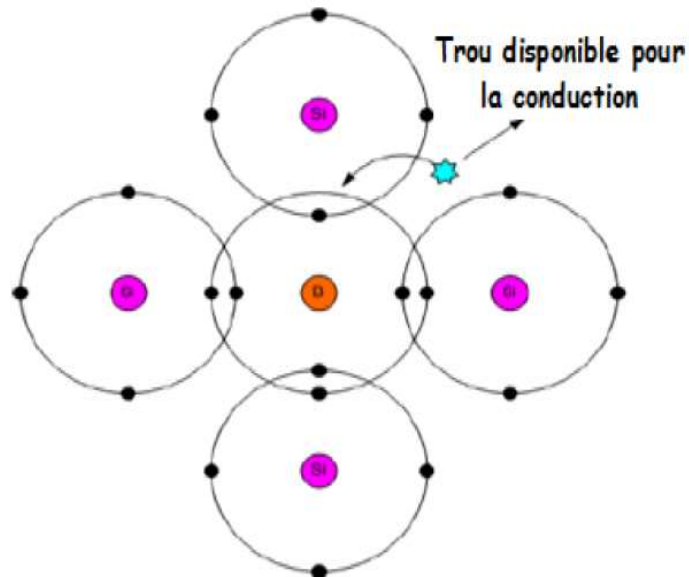
L'atome de bore

Le Bore (B) est le numéro 5 dans le tableau périodique



Il dispose de 5 protons et 5 électrons.

3 de ces électrons sont dans sa couche externe



	IIIA	IVA	VA	VIA
	5 B	6 C	7 N	8 O
	13 Al	14 Si	15 P	16 S
IIB	30 Zn	31 Ga	32 Ge	33 As
	48 Cd	49 In	50 Sn	51 Sb
		52 Te		

Silicium type-p (trous) dopants "accepteurs"

Figure 2.2 : un atome silicium, un atome de Bore (accepteur)

2.2.3. Semi-conducteur de type n

Supposons que nous ayons remplacé, dans un cristal de silicium, un atome de silicium par un atome d'arsenic. Ce dernier possède 5 électrons de valence alors que le silicium n'en possède que 4.

Cet électron supplémentaire est sur un niveau d'énergie, situé dans la bande interdite, placé à 54 meV en dessous du minimum, E_c , de la bande de conduction (**figure. 2.3**). C'est une orbitale localisée au voisinage de l'atome d'arsenic qui n'est pas délocalisée dans tout le cristal comme c'est le cas des bandes d'énergie. Mais, comme ce niveau d'énergie est situé à 54 meV au-dessous de la bande de conduction, l'électron de ce niveau passe très vite dans la bande de conduction par excitation thermique. Il participe donc à la conduction dans le cristal. Le dopage est qualifié pour cette raison de dopage **n**. Le dopant, dans cet exemple l'arsenic, se comporte comme un donneur d'électrons et le semi-conducteur obtenu est dit de type **n**.

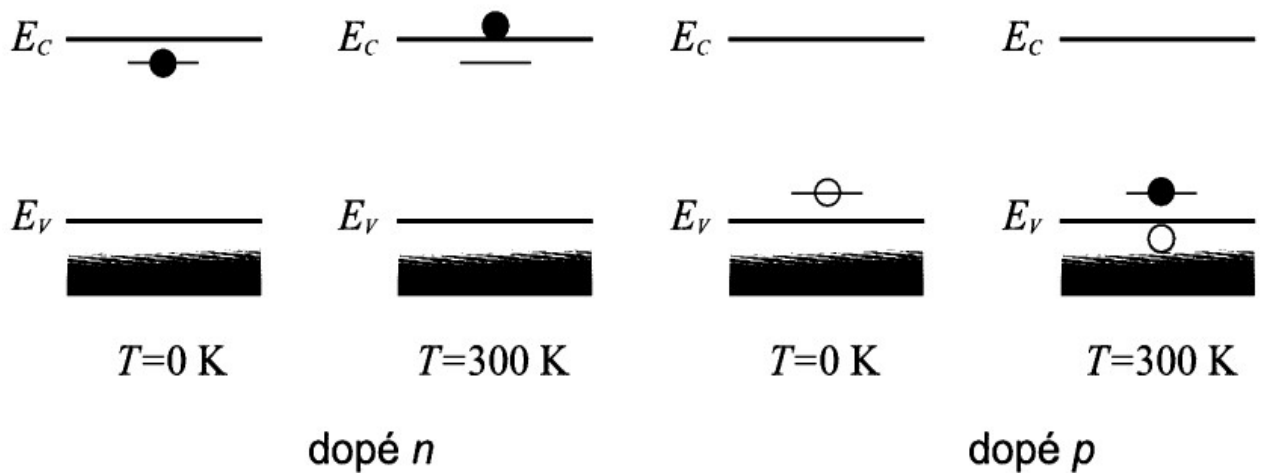


Figure 2.3. Structure de bande d'un semi-conducteur dopé n et dopé p à $T = 0\text{ K}$ et $T = 300\text{ K}$.
Un électron est représenté par un cercle plein et un trou par un cercle vide.

2.2.4. Semi-conducteur de type p

Considérons le silicium et dopons le avec des atomes de bore (élément de la colonne III). Comme le bore a un électron de valence de moins que le silicium, il se comporte comme un accepteur d'électrons. Une orbitale vide, localisée au voisinage de l'atome de bore se trouve à environ 50 meV au-dessus du maximum, E_v , de la bande de valence (**figure 2.3**). Compte tenu de la proximité de la bande de valence, un électron de celle-ci va très vite occuper ce niveau d'énergie par excitation thermique. Il en résulte un trou dans la bande de valence qui va contribuer au processus de conduction électrique. La création d'un

trou laisse un atome d'accepteur négatif, B^- , de charge $(-e)$, lié au réseau. Le dopage par des atomes accepteurs d'électrons est dit de type **p**.

2.2.5. Niveaux d'énergie des impuretés

Lors du dopage, un atome du réseau est substitué par un atome du dopant, donneur ou accepteur. À la position de substitution, et dans son voisinage, le potentiel électrique devient différent de ce qu'il était avant puisque l'atome est différent. On peut évaluer le niveau d'énergie que cela induit dans le schéma des bandes d'énergie si l'on fait l'hypothèse que la perturbation due à l'impureté est faible et de longue portée, i.e. qu'elle s'étend typiquement sur quelques dizaines d'atomes voisins. Avec cette hypothèse, on peut traiter l'atome de dopant par analogie avec l'atome d'hydrogène en supposant que la masse de l'électron est la masse effective au voisinage du bord de la bande d'énergie concernée. Nous allons, avec ces simplifications, traiter les dopants donneurs d'électrons puis ceux qui sont accepteurs.

a) Impuretés donneurs d'électrons

$$E_d = E_C - \left(\frac{1}{4\pi\epsilon} \right)^2 \frac{m_e^* e^4}{2\hbar^2} = E_C - 13,6 \left(\frac{m^*}{m} \right) \left(\frac{\epsilon_0}{\epsilon} \right)^2 \text{ eV}$$

Le rayon a de l'orbite de l'état fondamental de cet atome d'hydrogène un peu particulier peut être exprimé en fonction du rayon de Bohr de l'atome d'hydrogène a_0 par :

$$a = a_0 \left(\epsilon / \epsilon_0 \right) \left(m^* / m \right) = 0,53 \left(\epsilon / \epsilon_0 \right) \left(m^* / m \right)$$

b) Impuretés accepteurs d'électrons

$$E_a = E_V + \left(\frac{1}{4\pi\epsilon} \right)^2 \frac{e^4 m_h^*}{2\hbar^2} = E_V + 13,6 \left(\frac{m_h^*}{m} \right) \left(\frac{\epsilon_0}{\epsilon} \right)^2 \text{ eV}$$

Tableau 2.1 Niveau d'énergie des impuretés donneurs et accepteurs. Le niveau donneur est situé à E_d au-dessous du bas de la bande de conduction et le niveau accepteur à E_a au-dessus du haut de la bande de valence.

Semi-conducteur	Donneur	E_d (meV)	Accepteur	E_a (meV)
<i>AsGa</i>	<i>Si</i>	5,8	<i>Si</i>	35
<i>AsGa</i>	<i>Ge</i>	6,0	<i>Be</i>	28
<i>Si</i>	<i>As</i>	54	<i>B</i>	45
<i>Si</i>	<i>P</i>	45	<i>Ga</i>	72
<i>Ge</i>	<i>As</i>	13	<i>B</i>	10
<i>Ge</i>	<i>P</i>	12	<i>Ga</i>	11

2.2.6. Dopants amphotères et autodopage

Pour le silicium, par exemple, un dopant peut être accepteur ou donneur d'électrons mais pas les deux. Ce n'est pas le cas des composés **III-V** pour lesquels un même atome peut se comporter comme un donneur ou un accepteur selon sa position de substitution dans le cristal. Prenons le cas de l'AsGa, par exemple, avec le silicium comme dopant. Si l'atome de *Si* remplace un atome d'*As*, il se comporte comme un accepteur alors que s'il remplace un atome de *Ga* il se comporte comme une impureté donneur d'électrons. On le désigne pour cette raison sous le nom de dopant amphotère.

Pour certains semi-conducteurs composés de plusieurs types d'atomes, il peut y avoir autodopage. C'est le cas par exemple du composé *HgCdTe* pour lequel la substitution de sites occupés par des atomes de *Hg* ou de *Cd* par des atomes de *Te* peut conduire à des niveaux donneurs.

2.2.7. Niveaux profonds

Les éléments dopants que nous avons cités dans le cas du silicium conduisent à des niveaux d'énergie, situés dans la bande interdite, qui sont soit très proches de la bande de conduction pour un dopage n, soit très proches de la bande de valence pour un dopage de type p. Ils permettent d'enrichir le semi-conducteur en électrons (dopage n) ou en trous (dopage p). Il est également possible d'utiliser d'autres dopants que ceux cités plus haut, par exemple *Cu*, *Au*, *Ag*, *Fe*... Ils vont également conduire à la création d'un, ou parfois de deux, niveaux dans la bande interdite mais ceux-ci seront situés dans le milieu de celle-ci. On les appelle pour cette raison *des niveaux profonds*. Étant loin des bandes de conduction et de valence, ils ne vont pratiquement pas contribuer au phénomène de conduction électrique car l'énergie thermique n'est pas suffisante pour effectuer des transitions efficaces de ce niveau vers la bande de conduction, ou de la bande de valence vers ce niveau d'énergie. Par contre, leur présence va dégrader de manière appréciable les propriétés électriques du semi-conducteur en jouant le rôle de centres de recombinaison.

Les niveaux profonds peuvent se situer n'importe où dans la bande interdite et sont caractérisés par le fait qu'ils modifient localement la partie à courte portée du potentiel atomique puisque l'atome de substitution est très différent des autres atomes du cristal. La fonction d'onde associée à ces impuretés est très localisée dans l'espace alors que celle associée aux niveaux peu profonds, proches de la bande de conduction ou de la bande de valence, peut s'étendre sur plusieurs atomes voisins (le rayon de Bohr peut

atteindre la centaine d'Å). Des niveaux profonds peuvent également apparaître dans la bande interdite par suite de défauts dans le cristal, vacances ou dislocations, par exemple.

2.2.8. Dopage fort

Lorsque la concentration des dopants devient importante, typiquement supérieure à 10^{18} cm^{-3} , l'électron d'une impureté particulière peut être influencé par le potentiel créé par une impureté voisine (si deux impuretés sont situées à une distance inférieure à deux fois le rayon de Bohr introduit plus haut, i.e. typiquement 100 Å). La régularité de la présence d'impuretés peut induire des bandes d'énergie dans la bande interdite. La structure de bande du cristal hôte peut s'en trouver fortement perturbée et on observe une modification de la largeur de la bande interdite. C'est le cas pour le silicium dont le gap diminue pour des forts dopages.

2.3. Un semi-conducteur est dit « extrinsèque » ?

Les dopants de type n introduisent des niveaux d'énergie proches de la bande de conduction. Par conséquent, les électrons supplémentaires ainsi apportés passent pratiquement tous, à la température ordinaire, dans la bande de conduction. Pour les dopants de type p , il se forme aussi pratiquement autant de trous dans la bande de valence qu'il y a d'accepteurs supplémentaires.

Le but du dopage est d'augmenter soit le nombre d'électrons dans la bande de conduction, soit le nombre de trous dans la bande de valence. Ceci n'est réalisé que si les électrons situés sur les niveaux des dopants donneurs passent dans la bande de conduction, ou si les niveaux des impuretés accepteurs, situés juste au-dessus de la bande de valence, accueillent des électrons de cette dernière. Ceci n'a pas lieu au zéro absolu, ou à très basse température. Il faut une température suffisante pour que ce soit le cas.

Les porteurs de charge (électrons ou trous) provenant des dopants sont en général plus nombreux que ceux provenant du semi-conducteur lui-même. Dans un semiconducteur, le transport de charges se fait à la fois par les électrons et par les trous.

Dans un semi-conducteur intrinsèque, ces deux porteurs de charges sont en nombre égaux. Dans un semi-conducteur dopé, il y a un fort excédent d'électrons par rapport aux trous (dans le cas d'un semi-conducteur de type n) ou l'inverse (pour un semiconducteur de type p). On appelle porteurs

majoritaires, par opposition aux porteurs minoritaires, les porteurs de charges qui sont en plus grand nombre.

Pour un semi-conducteur de type n , appelons N_d le nombre d'atomes donneurs. Les électrons dans la bande de conduction ont deux origine :

- Les électrons provenant de la bande de valence, qui sont en nombre égal au nombre de trous n_h , puisqu'un électron passant de la bande de valence à la bande de conduction crée un trou dans la bande de valence.
- Les N_d électrons provenant du donneur.

Par conséquent $n_e = n_h + N_d$. Mais, puisque $N_d \gg n_h$, on a $n_e \approx N_d$. Comme $n_e n_h = n_i^2$ (loi d'action de masse où n_i^2 est la constante), $n_h = n_i^2 / N_d$. D'où, en résumé : $n_e = N_d$ et $n_h = n_i^2 / N_d$

Un raisonnement analogue pour un semi-conducteur de type p conduit à la relation: $n_h = n_e + N_a \approx N_a$

$$n_h = N_a \text{ et } n_e = n_i^2 / N_a$$

On peut ainsi calculer le niveau de Fermi d'un semi-conducteur extrinsèque : il est différent de celui du semi-conducteur intrinsèque. Nous le noterons pour le distinguer $E_F^{(i)}$.

2.3.1. Pour un semi-conducteur de type n :

$$n_e = 2 \left(\frac{m_e^* k_B T}{2\pi \hbar^2} \right)^{3/2} \exp \left(\frac{E_F - E_C}{k_B T} \right) = N_e \exp \left(\frac{E_F - E_C}{k_B T} \right)$$

Cette dernière expression est très générale et aucune hypothèse n'a été faite sur la nature du semi-conducteur. Elle exprime seulement la notion d'équilibre statistique. En utilisant l'expression précédente et $n_h = n_i^2 / N_d$, nous avons :

$$n_e = N_d = N_e \exp \left[- \left(E_C - E_F^{(n)} \right) / k_B T \right]$$

Où $E_F^{(n)}$ est le niveau de Fermi du semi-conducteur de type n . On déduit $E_C - E_F^{(n)} = K_B T \log(N_e / N_d)$

2.3.2. Pour un semi-conducteur intrinsèque

nous avons :

$$n_i = N_e \exp \left[- \left(E_c - E_F^{(i)} \right) / k_B T \right]$$

Où $E_F^{(i)}$ est le niveau de Fermi du semi-conducteur intrinsèque. Par conséquent :

$$E_F^{(n)} - E_F^{(i)} = -k_B T \left[\text{Log} \left(\frac{N_e}{N_d} \right) + \text{Log} \left(\frac{n_i}{N_e} \right) \right] = k_B T \text{Log} \left(\frac{N_d}{n_i} \right)$$

Comme $N_d \gg n_i$, cette différence est positive et $E_F^{(n)}$ est au-dessus de $E_F^{(i)}$. Le niveau de Fermi s'est rapproché de la bande de conduction (figure 2.4). Ce résultat était prévisible puisqu'en dopant le cristal on ajoute des particules au-dessus du niveau de Fermi du semi-conducteur intrinsèque.

Bande de conduction

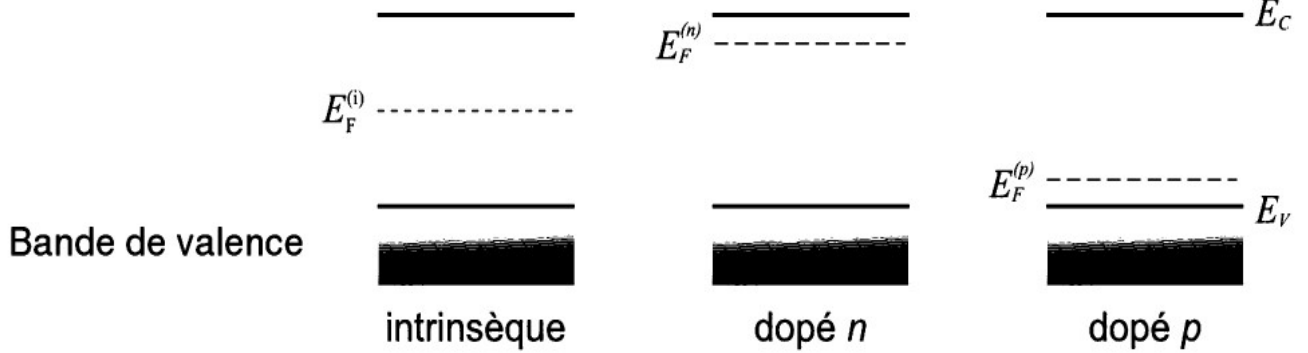


Figure 2.4. Position du niveau de Fermi selon la nature du semi-conducteur.

En suivant la même démarche pour un semi-conducteur de type p , on peut montrer que le niveau de Fermi d'un semi-conducteur dopé p est au-dessous de celui d'un semi-conducteur intrinsèque (figure 2.4). En effet, en utilisant l'expression :

$$n_h = 2 \left(\frac{m_h^* k_B T}{2\pi \hbar^2} \right)^{3/2} \exp \left(- \frac{E_F - E_v}{k_B T} \right) = N_h \exp \left(- \frac{E_F - E_v}{k_B T} \right)$$

on obtient :

$$E_F^{(p)} - E_F^{(i)} = -k_B T \text{Log} \left(\frac{N_a}{n_i} \right)$$

Chapitre 3.

Phénomène de Génération- Recombinaison

3.1. Etude des Semi-conducteurs hors équilibre

3.1.1. Définition d'un semi-conducteur hors équilibre

Dans un matériau semi-conducteur et en absence d'un champ électrique, le courant global créé par les électrons libres et les trous mobiles est nul. Lorsque le semi-conducteur est soumis à un champ électrique extérieur ($E \neq 0$) son comportement est totalement différent. En effet, le déplacement des électrons et des trous sous l'effet du champ électrique engendre un courant électrique.

Alors, sous l'effet du champ électrique

- L'électron libre subit une force électrique de forme ;

$$\vec{F} = -e \vec{E} = m_e \frac{d\vec{v}_e}{dt}$$

(m_e : est la masse effective de l'électron libre)

La vitesse de l'électron est donnée par ;

$$\vec{v}_e = -\frac{e \tau}{m_e} \vec{E} = -\mu_n \vec{E}$$

- μ_n est la mobilité électronique ; τ est le temps moyen entre deux collisions successives

3.1.2. Calcul des densités de courant

3.1.2.1. Calcul de la densité de courant de conduction (mobilité de porteurs de charges)

La densité de courant créée par les électrons est donnée par la relation :

$$\vec{J}_e = +n e \vec{v}_e$$

Avec n : la densité d'électrons

La densité de courant s'écrit aussi en fonction de la conductivité et le champ électrique :

$$\vec{J}_e = n e \mu_n \vec{E} = \sigma_e \vec{E}$$

Il résulte de ces deux relations :

$$\sigma_e = n e^2 \frac{\tau}{m_e} = n e \mu_n$$

- De manière analogue, la densité de courant créé par les trous mobiles s'écrit :

$$\vec{J}_h = p e \mu_p \vec{E} = \sigma_h \vec{E}$$

Avec :

$$\sigma_h = p e^2 \frac{\tau}{m_h} = p e \mu_p$$

m_h : est la masse effective du trou libre.

μ_p : est la mobilité des trous.

p : la densité de trous.

- La conductivité totale (σ) est la somme des deux conductivités, d'électrons (σ_e) et de trous (σ_h) :

$$\sigma = \sigma_e + \sigma_h$$

$$\sigma = \sigma_e + \sigma_h = n e \mu_n + p e \mu_p = e (n \mu_n + p \mu_p)$$

Dans le cas d'un semi-conducteur intrinsèque

$$(n = p = n_i)$$

La conductivité devient :

$$\sigma = \sigma_e + \sigma_h = e n_i (\mu_n + \mu_p)$$

3.1.3. Calcul de la densité de courant de diffusion

En absence du champ électrique ($E=0$), le courant total créé par le déplacement des porteurs de charges suivant les deux sens opposés (électrons et trous) est nul. On suppose que les densités de porteurs dans un volume délimité par l'intervalle $[-x, +x]$ sont constantes (figure 3.1). Dans le cas de la diffusion, la densité varie de part et d'autre de la surface (S) et sous agitation thermique les porteurs de charges se déplacent avec des vitesses variables.

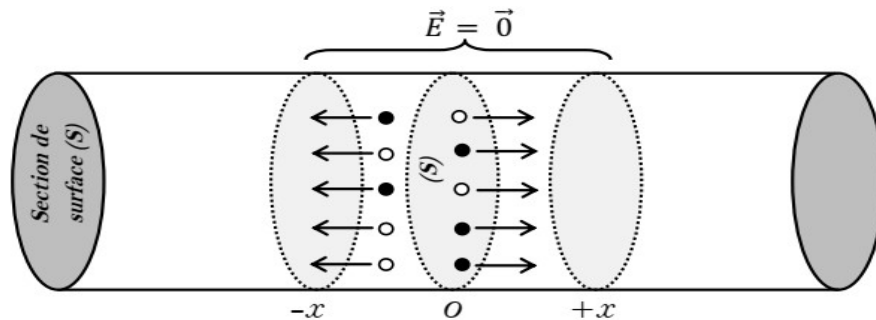


Figure 3.1 : Diffusion des porteurs de charges.

Pour simplifier l'étude de diffusion, nous considérons le cas d'une seule dimension, c'est-à-dire, des électrons se déplaçant suivant la direction (x) avec la même vitesse. Alors ;

La densité de courant de diffusion créé par les électrons en ($+x$) est donnée par la relation suivante :

$$J_{D1} = \frac{1}{2} n(+x) e v_{th} \quad \text{Avec : } +x = v_{th} \Delta t$$

La densité de courant de diffusion créé par les électrons en ($-x$) est donnée par la relation suivante :

$$J_{D2} = \frac{1}{2} n(-x) e v_{th} \quad \text{Avec : } -x = -v_{th} \Delta t$$

La densité totale J_{Dn} de courant de diffusion créé par les électrons est la somme des densités J_{D1} et J_{D2}

$$J_{Dn} = J_{D1} + J_{D2} = \frac{1}{2} n(+x) e v_{th} - \frac{1}{2} n(-x) e v_{th}$$

$$J_{Dn} = \frac{1}{2} e v_{th} [n(+x) - n(-x)]$$

Les densités d'électrons $n(+x)$ en ($+x$) et $n(-x)$ en ($-x$) peuvent être mises sous les formes suivantes :

$$n(+x) = \frac{dn}{dx} (+x) \quad \text{et} \quad n(-x) = \frac{dn}{dx} (-x)$$

La relation précédente de J_{Dn} s'écrit alors ;

$$J_{Dn} = \frac{1}{2} e v_{th} \left(\frac{dn}{dx} (+x) - \frac{dn}{dx} (-x) \right)$$

$$J_{Dn} = \frac{1}{2} e v_{th} \left(\frac{dn}{dx} x + \frac{dn}{dx} x \right) = \frac{1}{2} e v_{th} \left(2 x \frac{dn}{dx} \right)$$

$$J_{Dn} = \left(\frac{dn}{dx} \right) x e v_{th}$$

La vitesse thermique des électrons s'exprime en fonction de (x) et du temps de collision ($\Delta t = \tau$) par :

$$x = v_{th} \times \Delta t$$

Alors :

$$J_{Dn} = \frac{dn}{dx} (v_{th} \Delta t) e v_{th} = \frac{dn}{dx} (v_{th}^2 \tau) e$$

$$J_{Dn} = e v_{th}^2 \tau \frac{dn}{dx} = e D_n \frac{dn}{dx}$$

Avec la constante de diffusion des électrons.

$$D_n = v_{th}^2 \tau$$

De manière analogue, la densité de courant créé par les trous est donnée par la relation suivante :

$$J_{Dp} = -e v_{th}^2 \tau \frac{dp}{dx} = -e D_p \frac{dp}{dx}$$

Avec : D_p , la constante de diffusion des trous.

Les deux relations précédentes de densités de courant, appelé courant de diffusion (J_n et J_p), s'expriment en trois dimensions par les expressions vectorielles suivantes :

✓ Densité de courant de diffusion des électrons :

$$\vec{J}_{Dn} = e D_n \overrightarrow{\text{grad}}(n)$$

✓ Densité de courant de diffusion des trous :

$$\vec{J}_{Dp} = -e D_p \overrightarrow{\text{grad}}(p)$$

Enfin, la densité totale de courant de diffusion est la somme des densités de courant de diffusion créés par les électrons et les trous :

$$\vec{J}_D = \vec{J}_{Dn} + \vec{J}_{Dp} = e D_n \overrightarrow{\text{grad}}(n) - e D_p \overrightarrow{\text{grad}}(p)$$

En présence d'un champ électrique, les densités de courant sont :

➤ Densité de courant créé par les électrons.

$$\vec{J}_n = n e \mu_n \vec{E} + e D_n \overrightarrow{\text{grad}}(n)$$

➤ Densité de courant créé par les trous.

$$\vec{J}_p = p e \mu_p \vec{E} - e D_p \overrightarrow{\text{grad}}(p)$$

3.2. Phénomène de Génération-Recombinaison

Soient g' et r' respectivement le nombre de porteurs de charges créés par unité de volume et unité de temps ($cm^{-1}s^{-1}$) et le nombre de porteurs de charges qui disparaissent par unité de volume et de temps ($cm^{-1}s^{-1}$). Le nombre de porteurs de charges créés par unité de volume et unité de temps g' résulte d'une part de la génération spontanée due à l'agitation thermique g_{th} appelé taux de génération thermique et d'autre part de l'excitation par une source extérieure (g) telles que ; l'excitation optique, irradiation par particules, champ électrique ...etc

Le nombre de porteurs de charges créés par unité de volume et unité de temps g' s'écrit

$$g' = g + g_{th}$$

Le nombre de porteurs de charges r' est fonction des processus régissant la recombinaison des porteurs de charges excédentaires, c'est un paramètre propre au matériau.

La variation de la densité de porteurs de charges par unité de volume et unité de temps est due aux processus de génération – recombinaison produits sous l'effet de l'excitation extérieure et de l'agitation thermique, alors :

$$\left(\frac{dn}{dt}\right)_{gr} = g' - r' = g + g_{th} - r'$$

g' : c'est le taux de génération spécifique à l'excitation extérieure.

g_{th} et r' : paramètres spécifiques au matériau à une température donnée.

On pose :

$$r = r' - g_{th}$$

Cette relation représente le bilan entre la recombinaison et la génération thermique. r est un paramètre spécifique au matériau, il représente le taux de recombinaison.

La variation de la densité de porteurs de charges par unité de volume et unité de temps devient :

$$\left(\frac{dn}{dt}\right)_{gr} = g - r$$

3.2.1. Recombinaison directe (électron – trou)

Dans le processus de la recombinaison directe, le nombre d'électrons qui se recombinent est égal à celui de trous, soit ;

$$r'_n = r'_p = r' = k n p$$

Le taux de recombinaison r s'écrit donc :

$$r_n = r_p = r = r' - g_{th} = k n p - g_{th}$$

En absence de toute excitation extérieure ($g=0$), le taux de recombinaison r est nul (équilibre thermodynamique) et les densités (n) d'électrons et (p) de trous sont :

$$n = n_o, \quad p = p_o$$

Avec : $r=0$

La relation précédente devient :

$$r = k n_o p_o - g_{th} = 0 \Rightarrow g_{th} = k n_o p_o$$

D'où :

$$g_{th} = k n_i^2$$

Le taux de recombinaison s'écrit donc :

$$r = k (n p - n_i^2)$$

En régime hors équilibre les densités (n) d'électrons et (p) de trous peuvent être écrites sous forme :

$$n = (n_o + \Delta n) \quad \text{et} \quad p = (p_o + \Delta p)$$

Avec la condition de neutralité :

$$\Delta n = \Delta p$$

Le taux de recombinaison r s'écrit donc :

$$r = k [(n_o + \Delta n) (p_o + \Delta p) - n_i^2]$$

$$r = k [n_o \Delta p + p_o \Delta n + \Delta n \Delta p + n_o p_o - n_i^2]$$

Avec :

$$n_o p_o = n_i^2$$

$$r = k [n_o \Delta p + p_o \Delta n + \Delta n \Delta p]$$

Avec la condition de neutralité :

$$\Delta n = \Delta p$$

$$r = k [n_o + p_o + \Delta n] \Delta n = k [n_o + p_o + \Delta p] \Delta p$$

Le taux de recombinaison s'écrit aussi sous la forme suivante :

$$r = \frac{\Delta n}{\tau(\Delta n)} = \frac{\Delta p}{\tau(\Delta p)}$$

Avec :

$$\tau(\Delta n) = \frac{1}{k [n_o + p_o + \Delta n]}$$

$\tau(\Delta n)$: est la durée de vie des porteurs dans un semi-conducteur excité.

- Dans le cas d'un semi-conducteur de type (P) de faible injection, la densité intrinsèque n_o est très faible devant p_o

$$p_o \gg n_o \quad \text{et} \quad \Delta n = \Delta p \ll p_o$$

$$p = p_o + \Delta p \approx p_o \quad \text{et} \quad n = n_o + \Delta n$$

Il en résulte que :

$$\tau \approx \frac{1}{k p_o} \quad \text{et} \quad r \approx \frac{\Delta n}{\tau}$$

Dans le cas d'un semiconducteur de type (N) de faible injection, la densité intrinsèque p_o est très faible devant n_o .

$$n_o \gg p_o \quad \text{et} \quad \Delta n = \Delta p \ll n_o$$

$$p = p_o + \Delta p \quad \text{et} \quad n = n_o + \Delta n \approx n_o$$

Il en résulte que :

$$\tau \approx \frac{1}{k n_o} \quad \text{et} \quad r \approx \frac{\Delta p}{\tau}$$

Dans un matériau dopé en régime de faible injection, la densité de porteurs majoritaires est supposée constante. Le taux de recombinaison des porteurs minoritaires s'écrit :

- Pour un semi-conducteurs de type p : $r_n \approx \Delta n / \tau_n$

- Pour un semi-conducteurs de type n : $r_p \approx \Delta p / \tau_p$

Avec les durées de vie des porteurs minoritaires comme suit :

$$\tau_n = \frac{1}{k p_o} \quad \text{et} \quad \tau_p = \frac{1}{k n_o}$$

3.2.2. Recombinaison indirecte (centre de recombinaison)

Dans les semi-conducteurs, la durée de vie des porteurs dépend de leur densité. En effet, la probabilité pour qu'un électron et un trou se recombinent est très faible si les densités sont faibles. La présence d'impuretés a un effet très important sur la durée de vie des porteurs de charges, ils forment des centres de recombinaison. Il y a deux types de centres :

- Centre de recombinaison : capture d'un électron et d'un trou (recombinaison).
- Piège à électron : capture d'un électron puis le réémettre vers la bande de conduction

Le taux de recombinaison s'écrit :

$$r = \frac{1}{\tau_m} \left(\frac{p n - n_i^2}{2 n_i + p + n} \right)$$

$$\text{Avec : } \tau_m = \frac{1}{C N_R}$$

C : est le coefficient de capture, C_n (électrons) = C_p (trous) = C

N_R : Densité de centres de recombinaison.

3.2.2.1. Cas d'un semi-conducteur de type (N)

Dans ce cas d'une faible excitation :

$$n_o \gg n_i \gg p_o \text{ et } n = n_o + \Delta n \approx n_o$$

Et aussi :

$$p = p_o + \Delta p \ll n_o$$

Le taux de recombinaison donné par la relation précédente s'écrit :

$$r = \frac{1}{\tau_m} \left(\frac{p n - n_i^2}{n} \right) = \frac{p - n_i^2/n}{\tau_m} = \frac{p - p_o}{\tau_m}$$

$$r = \frac{\Delta p}{\tau_m}$$

3.2.2.2. Cas d'un semi-conducteur de type (P)

Dans ce cas d'une faible excitation :

$$p_o \gg n_i \gg n_o \text{ et } p = p_o + \Delta p \approx p_o$$

Et aussi :

$$n = n_o + \Delta n \ll p_o$$

Le taux de recombinaison donné par la relation précédente s'écrit :

$$r = \frac{1}{\tau_m} \left(\frac{p n - n_i^2}{p} \right) = \frac{n - n_i^2/p}{\tau_m} = \frac{n - n_o}{\tau_m}$$

$$r = \frac{\Delta n}{\tau_m}$$

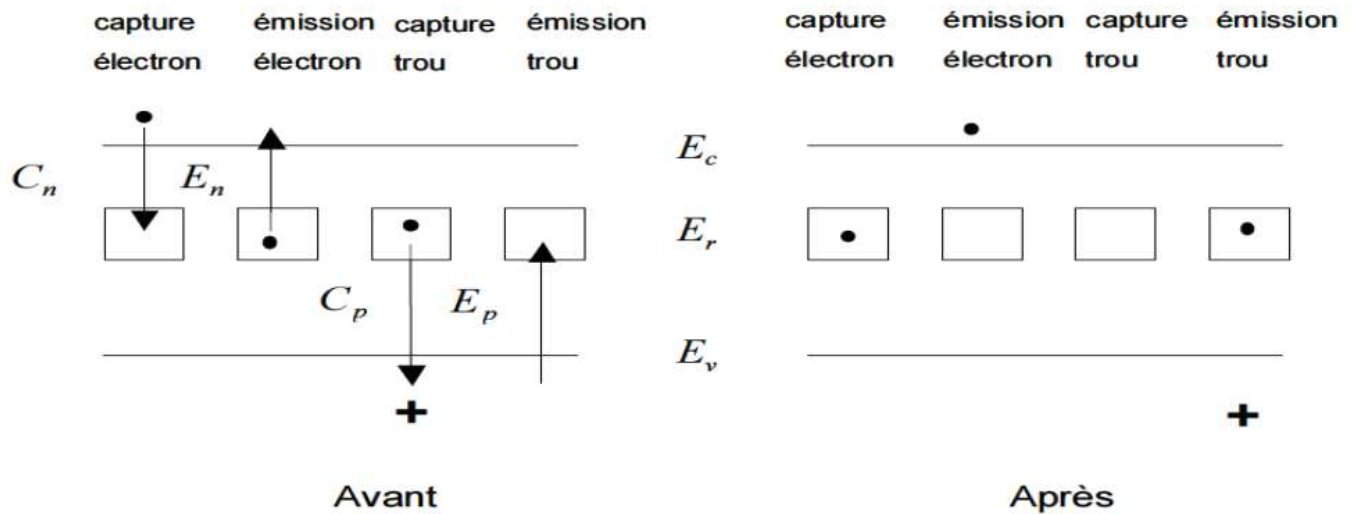


Figure 3.1 : Centres de recombinaisons

3.2.3. Zone de déplétion (dépeuplée)

Dans la zone dépeuplée les densités n et p sont négligeables devant n_i , alors :

$$n \ p \ll n_i^2$$

Le taux de recombinaison donné par la relation précédente devient :

$$r = - \frac{n_i}{2 \tau_m}$$

Le signe (-) indique que le nombre de porteurs créés thermiquement est plus important que le nombre de porteurs qui se recombinent. Un taux de recombinaison négatif correspond à une génération thermique de porteurs :

$$(n \ p \ll n_i^2).$$

Chapitre 4.

Etude des Jonctions PN

4.1. Définition des jonctions PN (abrupte et graduelle)

Une jonction PN est formée par la juxtaposition d'un semi-conducteur dopé type P (appelé anode) et d'un semi-conducteur dopé type N (appelé cathode), tous les deux d'un même monocristal semi-conducteurs, figure 4.1. Lorsque ces deux types de semi-conducteurs sont mis en contact, un régime électrique transitoire s'établit de part et d'autre de la jonction, suivi d'un régime permanent. Une jonction simple forme une diode.

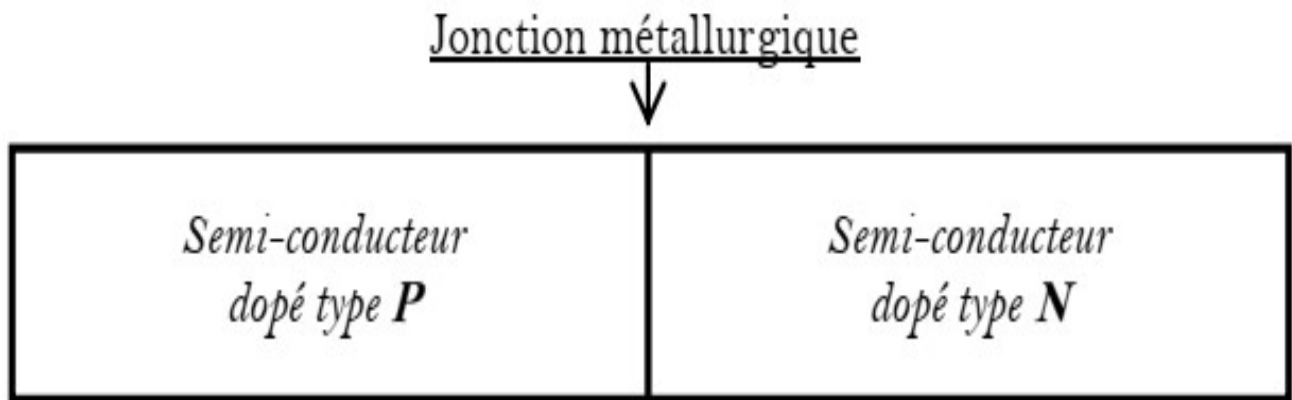


Figure 4.1 : Jonction métallurgique (PN)

4.1.1. Jonction PN abrupte :

Dans une jonction abrupte, la concentration en impuretés varie brutalement de la région dopée type P à la région dopée type N . C'est-à-dire, la différence $N_d - N_a$ passe d'une manière brutale à $x = 0$ d'une valeur négative dans la région dopée type P à une valeur positive dans la région dopée type N , voir figure 4.2a.

4.1.2. Jonction graduelle :

Dans une jonction graduelle, la concentration en impuretés est une fonction dépendante de x autour de la région de contact. C'est-à-dire, la différence $(N_d - N_a)$ dépend de x entre X_p et X_n , voir figure 4.2b, cas d'une dépendance linéaire.

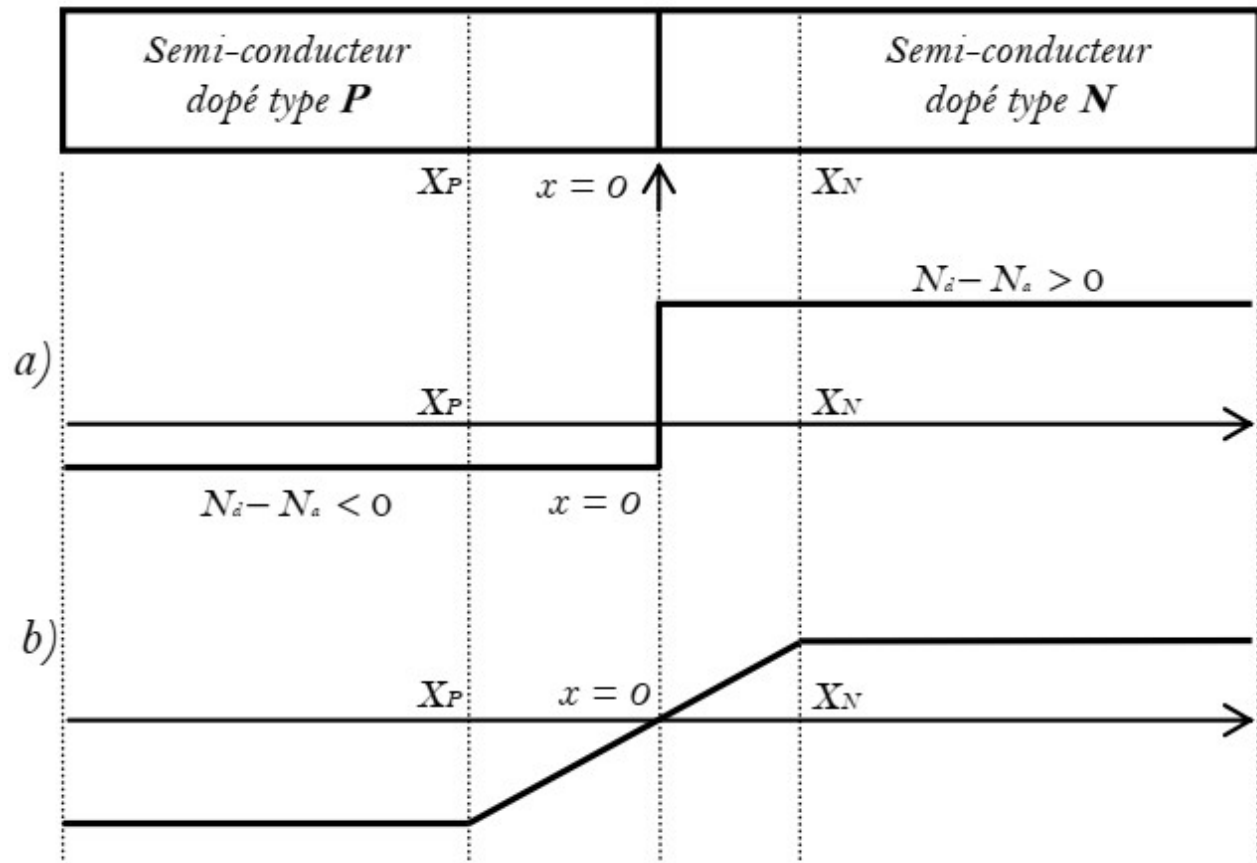


Figure 4.2 : Evolution de la différence ($N_d - N_a$). a) abrupte et b) graduelle

4.2. Etude d'une jonction PN abrupte non polarisée à l'équilibre

4.2.1. La charge d'espace $\rho(x)$: ($x_P < x < x_N$)

Considérons les deux types de semi-conducteurs avant la formation de la jonction, figure 4.3. Après la formation de la jonction entre les deux types de semi-conducteurs, un potentiel interne $\mathcal{O}_0$ de la jonction apparaît entre les deux niveaux de Fermi intrinsèques E_{Fi} . L'origine physique de ce potentiel interne $\mathcal{O}_0$ est la diffusion des porteurs de charges. En effet, la mise en contact des semi-conducteurs favorise la diffusion des électrons (e^-) de la région dopée N vers la région dopée P (pauvre en électrons) et laissent derrière eux des charges positives. De même, les charges positives (trous : h^+)

de la région dopée P diffusent vers la région dopée N (pauvre en trous) et laissent aussi derrière eux des charges négatives.

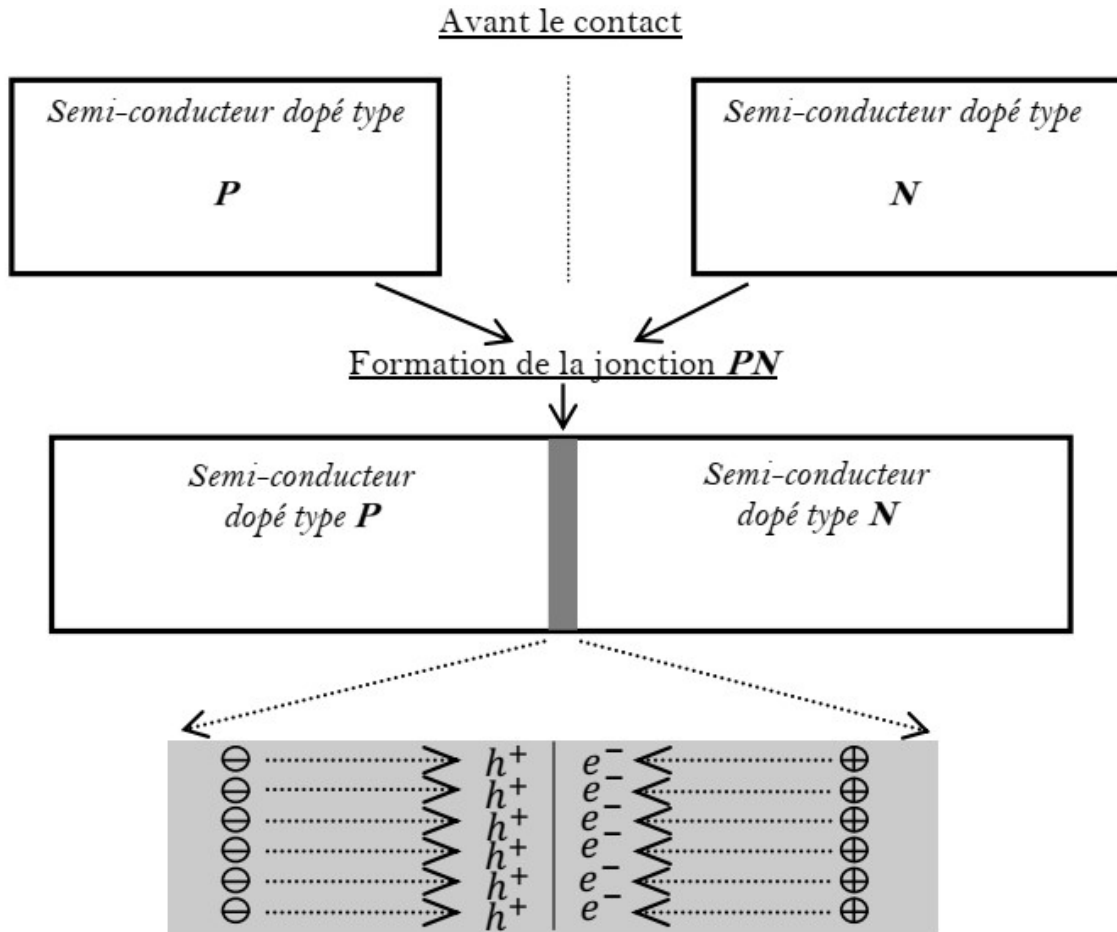


Figure 4.3 : Formation de la jonction (PN) et diffusion des porteurs de charges (électrons et trous)

A la fin de ce processus de diffusion des porteurs de charges, un équilibre permanent est établi et une zone pauvre en porteurs libres est formée. Cette zone est appelée **zone de charge d'espace** ou aussi « **zone de déplétion** ».

Afin d'étudier le potentiel électrostatique dans la jonction (zone de charge d'espace), la résolution de l'équation de Poisson est nécessaire. L'équation de Poisson s'écrit ;

$$\nabla^2 \phi(x, y, z) = \text{div} [\overrightarrow{\text{grad}} \phi(x, y, z)] = -\text{div} \vec{E} = -\frac{\rho(x)}{\epsilon_s} = -\frac{q}{\epsilon_s} (p - n + N_d^+ - N_a^-)$$

Avec :

$\rho(x)$: la densité de charge.

ϵ_s : La permittivité du semi-conducteur.

$$p(T) = n_i \exp\left(-\frac{E_{FP} - E_{Fi(P)}}{k_B T}\right) \exp\left(-\frac{q \phi(x)}{k_B T}\right)$$

$$n(T) = n_i \exp\left(-\frac{E_{FN} - E_{Fi(N)}}{k_B T}\right) \exp\left(\frac{q \phi(x)}{k_B T}\right)$$

Tel que :

$$N_d^+ = N_d$$

$$N_a^- = N_a$$

Cette équation peut être écrite sous la forme unidimensionnelle suivante :

$$\frac{d^2 \phi}{dx^2} = -\frac{\rho(x)}{\epsilon_s} = -\frac{q}{\epsilon_s} (p - n + N_d^+ - N_a^-)$$

Pour résoudre cette équation, nous supposons que la charge présente dans le semi-conducteur est seulement due à une distribution homogène d'impuretés. La concentration en porteurs de charges libres est donc négligeable devant N_d et N_a . En plus, la densité de charge est supposée constante dans les deux régions de la zone de déplétion (cas d'une jonction abrupte), figure 4.4.

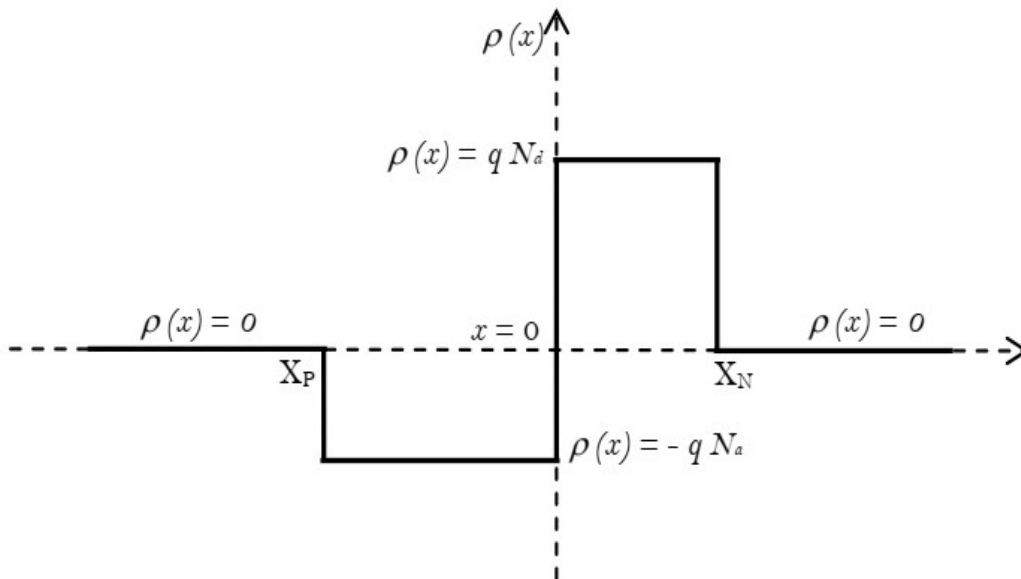


Figure 4.4 : Evolution de la densité de charge $\rho(x)$.

Alors, la distribution de charges $\rho(x)$ est donnée par la relation suivante :

$$\rho(x) = \begin{cases} 0, & \text{si } -\infty < x \leq -X_p \\ -q N_a, & \text{si } -X_p < x \leq 0 \\ +q N_d, & \text{si } 0 \leq x < +X_N \\ 0, & \text{si } +X_N \leq x < +\infty \end{cases}$$

Avec : $q = e$, la charge de l'électron

4.2.2. Calcul des champs électriques : $E_P(x)$ et $E_N(x)$.

Pour déterminer le champ électrique créé dans chaque région il faut intégrer l'équation de Poisson correspondante. Figure 4.5.

- $-\infty < x \leq X_p$ avec : $\rho(x)=0$, le champ électrique est nul ($E=0$).
- $X_p < x \leq 0$ avec : $\rho(x) = -qN_a$, L'équation de Poisson s'écrit :

$$\frac{d^2 \phi}{dx^2} = - \frac{dE(x)}{dx} = - \frac{\rho(x)}{\epsilon_s} = \frac{q N_a}{\epsilon_s}$$

Le champ électrique au point $x = X_p$ est nul, alors par intégration de l'équation de Poisson le champ électrique dans la région dopée P de la jonction s'écrit:

$$E_P(x) = - \frac{q N_a}{\epsilon_s} (x - X_p)$$

- $0 < x \leq +X_N$ avec : $\rho(x) = qN_d$

L'équation de Poisson s'écrit :

$$\frac{d^2 \phi}{dx^2} = - \frac{dE(x)}{dx} = - \frac{\rho(x)}{\epsilon_s} = - \frac{q N_d}{\epsilon_s}$$

Le champ électrique au point $x = +X_N$ est nul, alors par intégration de l'équation de Poisson, le champ électrique dans la région dopée N de la jonction s'écrit:

$$E_N(x) = - \frac{q N_d}{\epsilon_s} (X_N - x)$$

Au point ($x=0$), la continuité du champ électrique impose que :

$$\frac{q N_a}{\epsilon_s} X_P = \frac{q N_d}{\epsilon_s} X_N \Rightarrow N_a X_P = N_d X_N$$

- $+X_N < x \leq +\infty$ avec : $\rho(x)=0$, le champ électrique est nul ($E=0$).

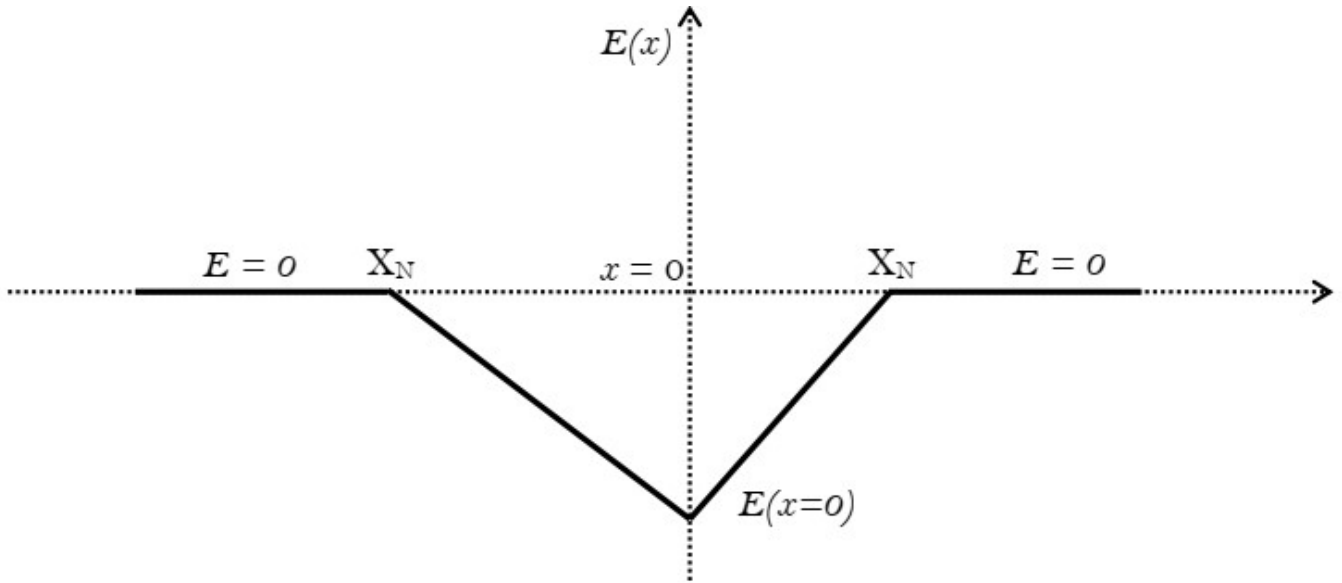


Figure 4.5 : Champ électrique dans les différentes régions de la jonction.

4.2.3. Calcul des potentiels électriques : $V_P(x)$ et $V_N(x)$

Le potentiel électrique ϕ créé dans chaque région de la jonction est déterminé par intégration du champ électrique E .

$$E(x) = -\frac{d\phi}{dx} = -\frac{dV}{dx} \Rightarrow V = -\int E(x) dx$$

- $-\infty < x \leq X_P$ avec : $\rho(x)=0$, le champ électrique est nul ($E=0$), le potentiel électrique est constant.

Alors : $V=V_P(x=X_P)$

- $X_P < x \leq 0$: En $x=X_P$ ($E=0$ et $V=V_P$)

Le champ électrique E est donné par la relation précédente :

$$E_P(x) = -\frac{q N_a}{\epsilon_s} (x - X_P)$$

$$E(x) = -\frac{dV}{dx} \Rightarrow \int_{V_P}^{V(x)} dV = -\int_{X_P}^x E(x) dx$$

$$V(x) = \frac{q N_a}{\epsilon_s} \int_{X_P}^x (x - X_P) dx = \frac{q N_a}{2 \epsilon_s} (x - X_P)^2 + V_P$$

$$V_P(x) = \frac{q N_a}{2 \epsilon_s} (x - X_P)^2 + V_P$$

- $0 \leq x < +X_N$: En $x=X_N$ ($E=0$ et $V=V_N$)

Le champ électrique E est donné par la relation précédente :

$$E_N(x) = -\frac{q N_d}{\epsilon_s} (X_N - x)$$

$$E(x) = -\frac{dV}{dx} \Rightarrow \int_{V(x)}^{V_N} dV = -\int_x^{X_N} E(x) dx$$

$$V(x) = +\frac{q N_d}{\epsilon_s} \int_x^{X_N} (X_N - x) dx = -\frac{q N_d}{\epsilon_s} \int_x^{X_N} (x - X_N) dx = -\frac{q N_d}{2 \epsilon_s} (x - X_N)^2 + V_N$$

$$V_N(x) = -\frac{q N_d}{2 \epsilon_s} (x - X_N)^2 + V_N$$

Au point, la continuité du potentiel impose que :

$$\frac{q N_a}{2 \epsilon_s} X_P^2 + V_P = -\frac{q N_d}{2 \epsilon_s} X_N^2 + V_N$$

$$V_N - V_P = \frac{q N_a}{2 \epsilon_s} X_P^2 + \frac{q N_d}{2 \epsilon_s} X_N^2 = \frac{q}{2 \epsilon_s} (N_a X_P^2 + N_d X_N^2)$$

Le potentiel de diffusion V_d est la différence entre le potentiel V_N et V_P , il s'écrit ;

$$V_d = V_N - V_P = \frac{q N_a}{2 \epsilon_s} X_P^2 + \frac{q N_d}{2 \epsilon_s} X_N^2 = \frac{q}{2 \epsilon_s} (N_a X_P^2 + N_d X_N^2)$$

$$V_d = V_N - V_P = \frac{q}{2 \epsilon_s} (N_a X_P^2 + N_d X_N^2)$$

- $+X_N \leq x < +\infty$: le champ électrique E est nul ($E=0$), Le potentiel est donc constant.

Alors ; $x=X_N$, $E=0$ et $V=V_N$ ($x=X_N$)

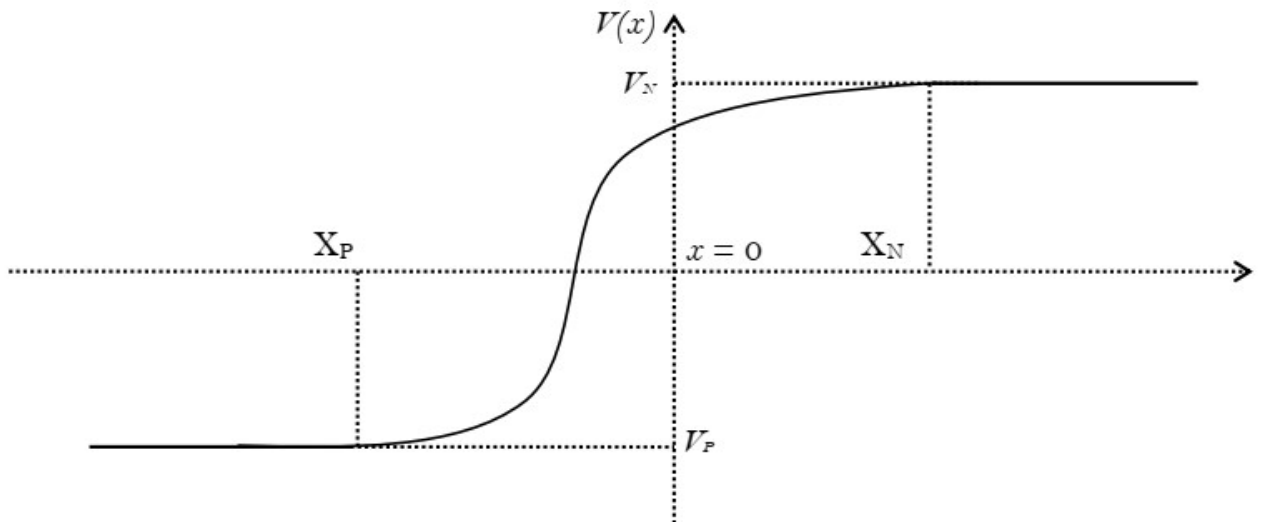


Figure 4.6 : Potentiel électrique dans les différentes régions de la jonction

4.2.4. Calcul de la tension de diffusion V_d ; barrière de potentiel.

Dans une jonction abrupte, la concentration en atomes donneurs dans la région dopée N est N_d (en cm^{-3}) et la concentration en accepteurs dans la région dopée P est N_a (en cm^{-3}), la différence en énergie entre le niveau de Fermi du semi-conducteur intrinsèque (E_{Fi}) et le niveau de Fermi du semi-conducteur extrinsèque (E_{FN}) ou (E_{FP}) s'écrit : figure 4.7.

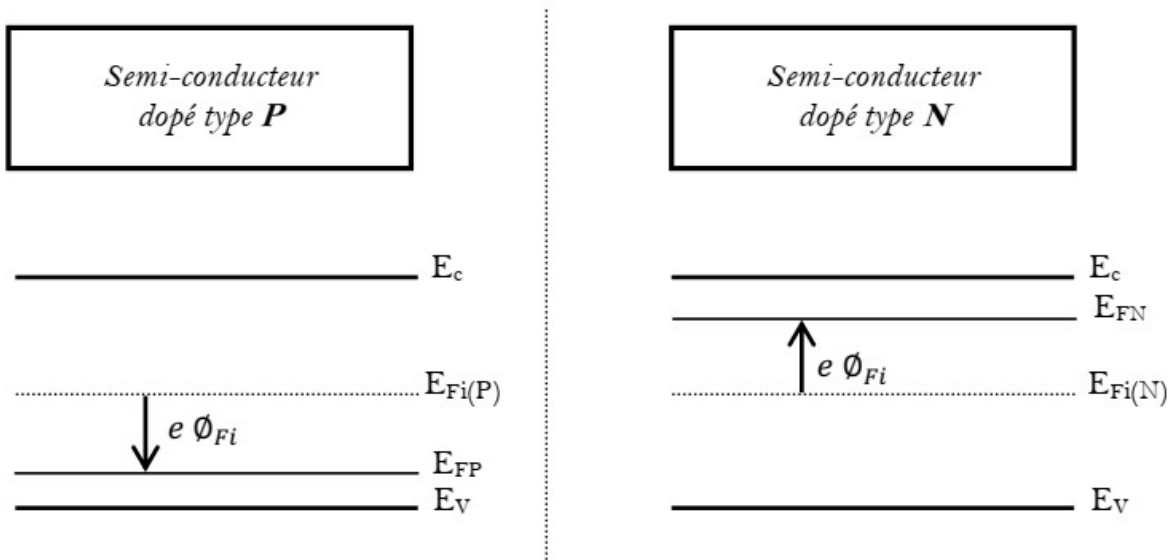


Figure 4.7 : Diagrammes énergétiques des semi-conducteurs (dopé type P et dopé type N) avant la formation de la jonction.

- Cas d'un semi-conducteur dopé N :

$$E_{FN} - E_{Fi(N)} = e \phi_{Fi(N)} = k_B T \ln \left(\frac{N_d}{n_i} \right)$$

Avec ;

$$n_N = N_d \text{ et } p_N = n_i^2 / N_d$$

- Cas d'un semi-conducteur dopé P :

$$E_{Fi(P)} - E_{FP} = e \phi_{Fi(P)} = k_B T \ln \left(\frac{N_a}{n_i} \right)$$

Avec ;

$$p_P = N_a \text{ et } n_P = n_i^2 / N_a$$

L'énergie ($e \phi_0$) correspondante au potentiel interne est déduite de la condition d'alignement des deux niveaux de Fermi extrinsèques ($E_{FN} = E_{FP}$), figure 4.8.

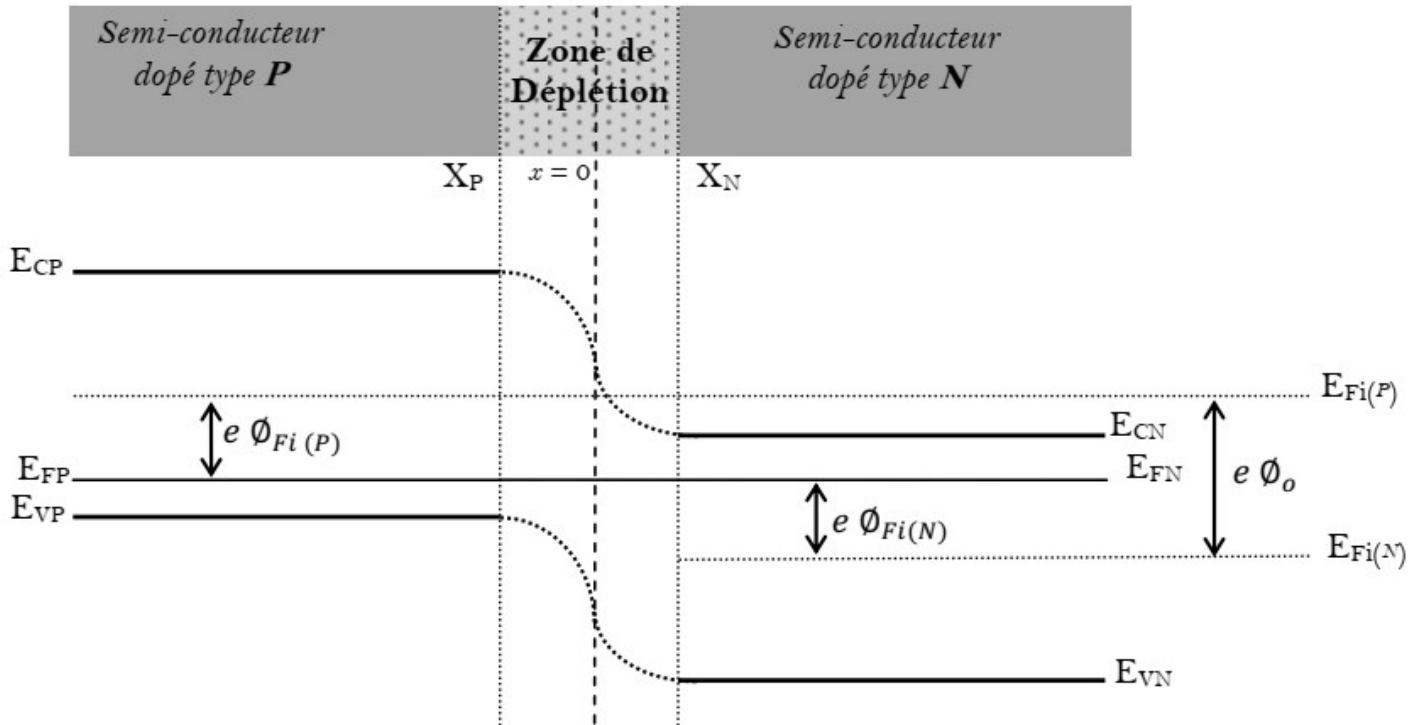


Figure 4.8 : Diagramme énergétique des semi-conducteurs (dopé type P et dopé type N) après la formation de la jonction

Le potentiel interne s'écrit :

$$E_{Fi(P)} - E_{Fi(N)} = e (\phi_{Fi(N)} + \phi_{Fi(P)}) = e \phi_o = k_B T \ln\left(\frac{N_d}{n_i}\right) + k_B T \ln\left(\frac{N_a}{n_i}\right)$$

$$e \phi_o = k_B T \ln\left(\frac{N_d N_a}{n_i^2}\right)$$

$$V_d = \phi_o = \frac{k_B T}{e} \ln\left(\frac{N_d N_a}{n_i^2}\right)$$

Le potentiel de diffusion V_d s'écrit aussi ;

$$V_d = \frac{E_{CP} - E_{CN}}{e}$$

A température ambiante, la tension de diffusion de V_d est de l'ordre de 0,7 V pour une jonction au silicium est de l'ordre de 0,35 V pour une jonction germanium.

4.2.5. Calcul de la largeur de la zone de charge d'espace w (zone de déplétion).

Les relations précédentes du potentiel de diffusion permettent de déduire la largeur de la zone de déplétion

$$w = X_P + X_N$$

✓ La largeur de la région dopée P est donnée par la relation :

$$X_P = \sqrt{\frac{2 \varepsilon_s}{e} \frac{V_d N_d}{N_a (N_a + N_d)}}$$

✓ La largeur de la région dopée N est donnée par la relation :

$$X_N = \sqrt{\frac{2 \varepsilon_s}{e} \frac{V_d N_a}{N_d (N_a + N_d)}}$$

✓ La largeur w est donnée par la relation :

$$w = \sqrt{\frac{2 \varepsilon_s}{e} \frac{V_d (N_a + N_d)}{N_a N_d}}$$

4.3. Etude d'une jonction PN polarisée (hors équilibre)

4.3.1. Polarisation directe

Dans le cas d'une polarisation directe le sens passant d'une diode est défini par le sens des courants créés par les porteurs majoritaires dans la jonction. Les électrons majoritaires de la zone N se déplacent vers l'anode, les trous majoritaires de la zone P se déplacent vers la cathode. Le sens direct est défini par un courant dirigé de l'anode vers la cathode, figure 4.9.

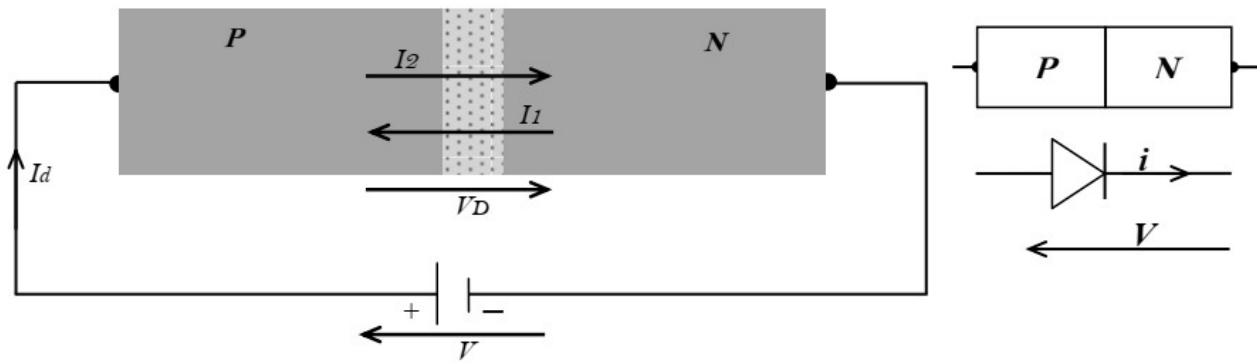


Figure 4.9: Polarisation directe de la jonction PN et symbole de la diode

La différence de potentiel est supérieure à zéro.

$$V = V_p - V_n > 0$$

Le courant crée par les porteurs minoritaires s'écrit :

$$|i_1| = I_s = i_o e^{\frac{-e V_D}{k_B T}}$$

Et le courant des porteurs majoritaires s'écrit :

$$|i_2| = I_o e^{\frac{-(e V_D - V)}{k_B T}} = I_s e^{\frac{e V}{k_B T}} \gg |i_1|$$

Le courant direct s'écrit donc :

$$i_d = |i_2| - |i_1| = I_s \left(e^{\frac{e V}{k_B T}} - 1 \right) \approx I_s e^{\frac{e V}{k_B T}} \gg |i_1|$$

Alors, le courant i_d résulte du déplacement des porteurs majoritaires

En polarisation directe, il faut vérifier que l'intensité reste inférieure à une valeur maximale $i_{\max}$ qui peut varier de 20 mA pour une diode de signal utilisable en hautes fréquences à plusieurs ampères pour une diode de redressement utilisable à une fréquence voisine à 50 Hz.

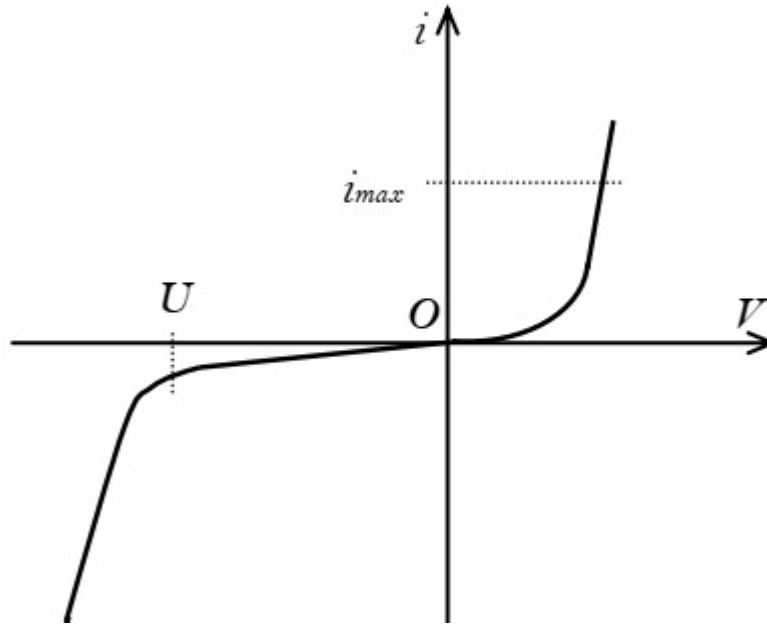


Figure 4.9: Caractéristique $I = f(V)$ d'une diode

4.3.2. Polarisation inverse.

Lors de l'utilisation en polarisation inverse, la tension doit rester supérieure à une valeur minimale notée $(-U)$. Dans le cas de la polarisation inverse la jonction est reliée à une alimentation dans le sens bloqué, figure 4.10.

La différence de potentiel s'écrit :

$$V = V_n - V_p > 0$$

✓ Le courant crée par les porteurs minoritaires s'écrit :

$$|i_1| = I_s = i_o e^{\frac{-eV_D}{k_B T}}$$

✓ Et le courant des porteurs majoritaires s'écrit :

$$|i_2| = I_o e^{\frac{-(eV_D + V)}{k_B T}} = I_s e^{\frac{-eV}{k_B T}} \ll |i_1|$$

✓ Le courant inverse s'écrit donc :

$$|i_{inv}| = |i_1| - |i_2| = I_s \left(1 - e^{\frac{-eV}{k_B T}} \right) \approx I_s$$

Alors, le courant i_{inv} résulte du déplacement des porteurs minoritaires. L'intensité du courant est très faible, il est de l'ordre de quelques pA à quelques nA

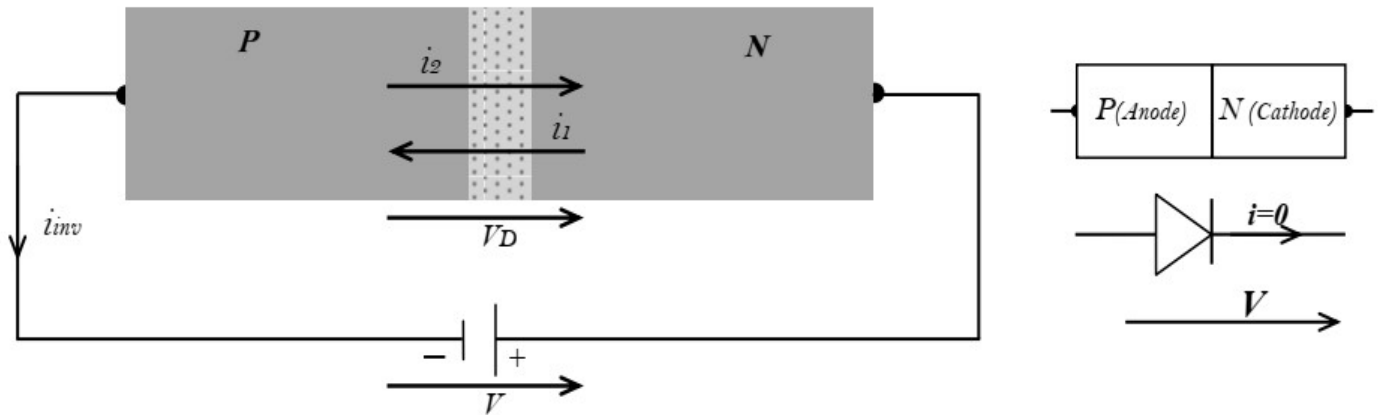


Figure 4.10: Polarisation inverse de la jonction PN

4.3.2. Caractéristique courant – tension $I(V)$

La caractéristique courant - tension $i(V)$ est représentée dans la figure 4.11 suivant :

✓ La résistance statique est donnée par la relation :

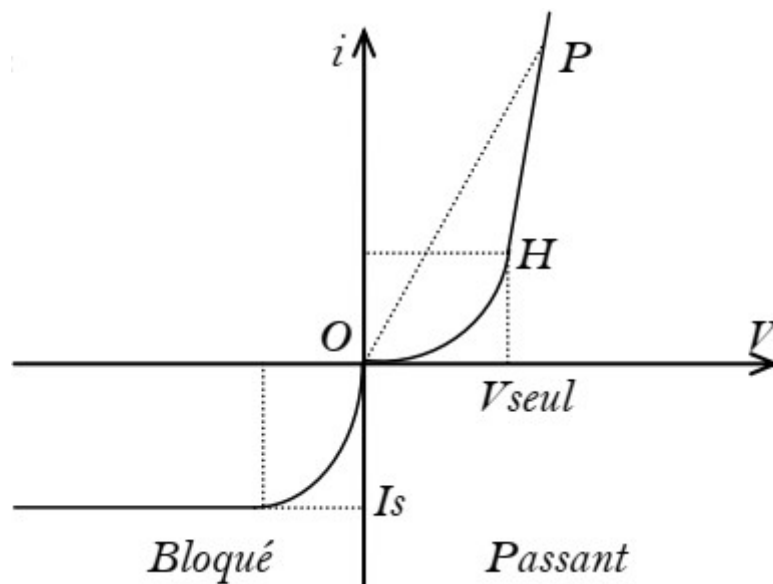


Figure 4.11: Caractéristique $I = f(V)$ d'une diode.

Remarques :

- ✓ La résistance dans le sens passant R_a est de quelques Ω .
- ✓ La résistance dans le sens bloqué est très élevée et tend vers infinie.
- ✓ Le courant inverse I_s est constant et vaut quelques μA .
- ✓ La tension seuil $V_{seuil} = V_D = 0,2 \text{ V}$ pour Ge et $0,6 \text{ V}$ pour Si.
- ✓ La zone où le potentiel est inférieur à $-U$ ($u < -U$), est appelée zone d'avalanche. C'est une zone dans laquelle la diode est détériorée.

4.4. Types de jonctions PN**4.4.1. Diodes Zener**

Les diodes Zener sont des stabilisateurs de tensions continues. Ce type de diode permet de conserver la tension constante dans la zone de claquage. La caractéristique de la diode Zener est représentée par la fonction $I = f(V)$ suivante :

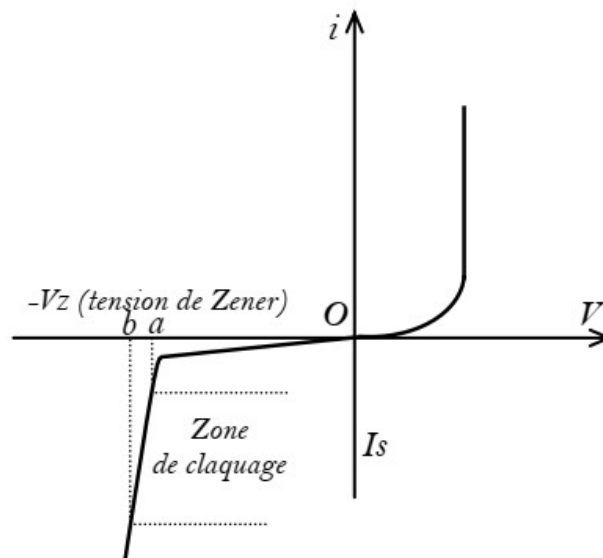
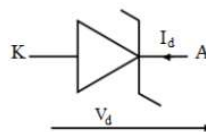


Figure 4.13 : Caractéristique $I = f(V)$ d'une diode Zener



Symbole d'une diode Zener

En polarisation inverse, si la tension dépasse une tension ($-V_Z$), selon de composant choisi, il apparaît un courant dit de claquage de la jonction. Ce phénomène est dû soit à l'effet d'avalanche, soit à l'effet Zener. La tension de claquage est faible (quelques volts ou quelques dizaines de volts).

- ✓ Le claquage par effet Zener se produit lorsque

$$|V| < a.$$

- ✓ Le claquage par l'effet avalanche se produit lorsque

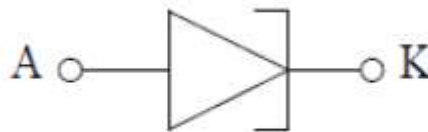
$$|V| > b$$

4.4.2. Diodes à avalanche

Dans les diodes à semi-conducteurs et même les transistors, le mode de claquage le plus courant se produit par effet avalanche. Ce phénomène résulte lorsqu'une forte tension inverse est appliquée aux bornes de la jonction. En effet, sous l'effet du champ électrique interne, l'énergie cinétique des porteurs minoritaires devient suffisante pour créer des paires électron-trou dans la zone de transition. Ces porteurs sont ensuite accélérés par le champ interne et créent à leur tour de nouvelles paires électron-trou (d'où le nom d'avalanche du phénomène). Le courant augmente très rapidement et provoque ainsi la destruction de la jonction par effet Joule.

4.4.3. Diode à effet tunnel

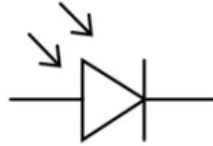
La diode à effet tunnel désigne une diode dont les zones N et P sont hyper dopées. Ainsi, en polarisation inverse, les électrons de la bande de valence de la zone P peuvent passer directement à la bande de conduction de la zone N. Le passage de ces porteurs, par effet tunnel, entraîne l'apparition d'un courant dû au franchissement de la barrière de potentiel dans la zone de charge d'espace. Ce processus de nature quantique est appelé effet tunnel.



Symbole d'une diode à effet tunnel.

4.4.5. Photodiode

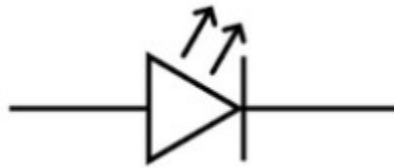
Une photodiode est un composant constitué d'une jonction PN. Elle a la capacité de détecter un rayonnement du domaine optique et de le transformer en signal électrique.



Symbole d'une photodiode

4.4.6. Diode électroluminescente

Une diode électroluminescente est un composant formé d'une jonction PN polarisée en direct. Suite à des recombinaisons radiatives entre les porteurs minoritaires qui se retrouvent en excès par rapport à l'équilibre, on obtient un émetteur optique capable d'émettre un rayonnement.



Symbole d'une diode électroluminescente.

4.4.7. Diode Laser

Les diodes lasers sont des émetteurs de rayonnement tout comme les diodes électroluminescentes, mais le rayonnement émis aura des propriétés différentes. Il sera davantage monochromatique et cohérent (puissance optique).

4.4.8. Cellules solaires

Une cellule solaire est un composant électronique, constitué d'un matériau semiconducteur, qui convertit la lumière du soleil en électricité. Le principe de la conversion photovoltaïque est basé sur l'absorption des photons incidents et la création de paires électron-trou, si l'énergie du photon incident est supérieure au gap du matériau. Une partie des électrons ne revient pas à son état initial et les électrons arrachés de la bande de valence vont créer une tension électrique continue et faible.

4.5. Applications des jonctions PN :

4.5.1. Redressement de signaux alternatifs

Les appareils électroniques fonctionnent sous tension continue. Pour cela, la diode à jonction est principalement utilisée pour transformer un signal sinusoïdal en signal continu. Parmi les applications des diodes les plus utilisés sont le redressement mono-alternance et double alternance.

- ✓ Le redressement mono-alternance consiste à transformer le signal sinusoïdal, c'est-à-dire à supprimer les alternances négatives ou les alternances positives, figure 4.13.
- ✓ Le redressement double alternance consiste à rendre positive les alternances négatives du signal sinusoïdal, figure 4.13.

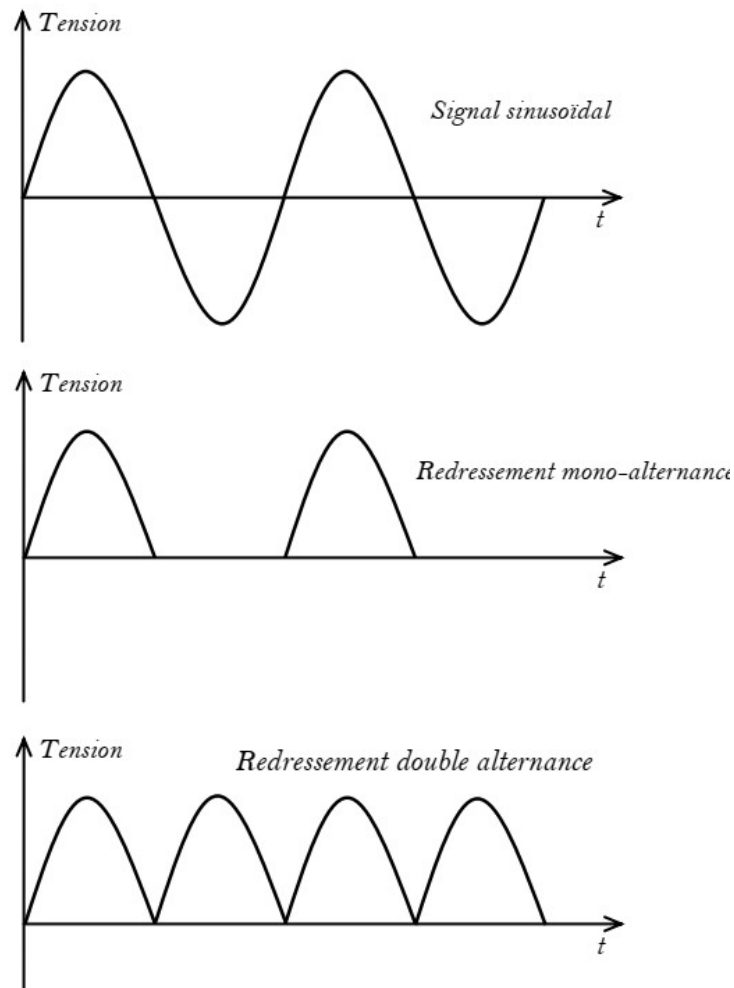


Figure 4.13: Redressement mono-alternance et double alternance d'un signal sinusoïdal.

4.5.2. Commutation

On appelle commutation le passage de la jonction d'un état à l'autre (par exemple du régime inverse : état bloqué au régime direct : état passant). Dans ces conditions, la capacité de transition intervient lorsque la jonction est en état bloqué, et lorsque la jonction est en état passant, c'est la capacité de diffusion qui intervient. Le temps nécessaire pour le passage d'un état bloqué à un état passant est appelé : temps de commutation.

Chapitre 5.

Dispositifs à hétérojonction

5.1. Jonction métal-métal

Nous allons étudier la jonction de deux métaux en contact qui est une bonne introduction à la jonction entre un semi-conducteur et un métal. Nous avons défini plus haut l'énergie minimum qu'il faut fournir pour arracher un électron à un métal ; c'est le **travail** (ou **potentiel**) **d'extraction** W . Il représente la différence entre l'énergie du vide et celle du niveau de Fermi, E_F (figure 5.1 (a)).

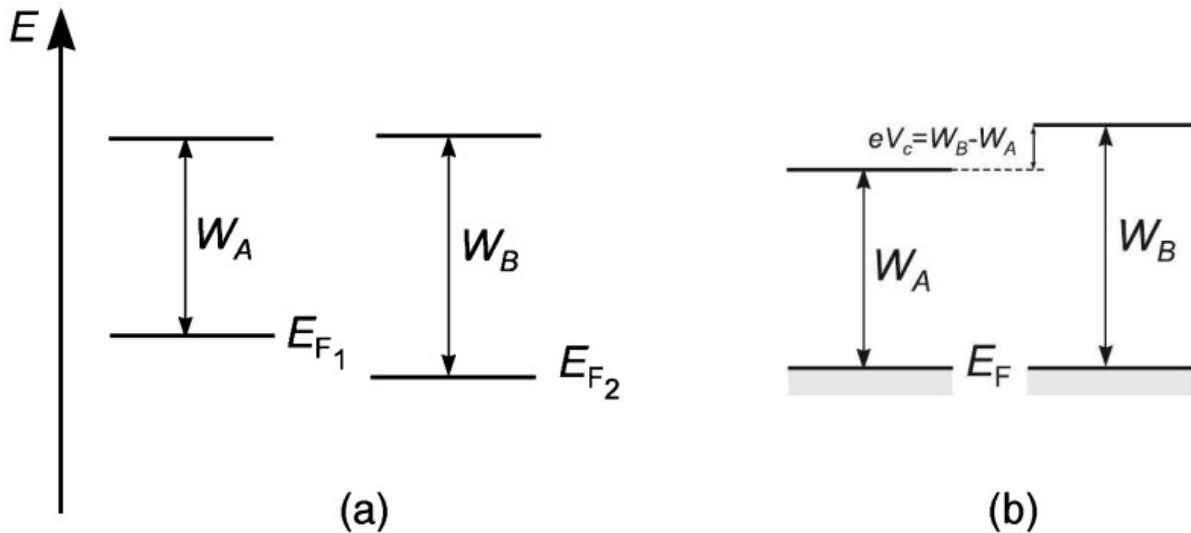


Figure 5.1 : Principe de formation de la jonction métal-métal. La figure (a) montre le diagramme d'énergie des deux métaux séparés. La figure (b) correspond au cas où la jonction métal-métal est formée.

Considérons deux métaux A et B ayant, respectivement, des niveaux de Fermi E_{F1} et E_{F2} , et des travaux d'extraction W_A et W_B (figure.5.1 (a)). Nous supposons $E_{F1} > E_{F2}$ pour la discussion. Lors de la création de la jonction métal-métal, les niveaux de Fermi s'équilibrent. Pour cela, des électrons du métal **B** vont passer vers le métal **A**. Il s'ensuit que le voisinage de la jonction se charge positivement du côté du métal **B** car les atomes du réseau sont fixes et que des électrons ont quitté cette région.

Au contraire, le voisinage de la jonction du métal **A** se charge négativement par suite d'un excès d'électrons. À l'équilibre, *i.e.* lorsque les potentiels chimiques des deux métaux sont égaux, il y a un champ électrique, qui empêche que tous les électrons de **B** passent en **A**. Le potentiel de contact V_c vaut

$V_c = (W_B - W_A) / e$. Ce potentiel de contact est indiqué sur la figure 5.1 (b). Pour un contact entre de l'argent et de l'or, par exemple, on a $V_c = +0,5$ V ($W_{Ag} = 4,33$ eV et $W_{Au} = 4,83$ eV).

Le thermocouple est l'application la plus simple de la jonction métal-métal. On peut aussi considérer des configurations de type A-B-A. Si l'on fait circuler un courant dans ce dispositif, on observe que la température d'une des jonctions augmente alors que celle de l'autre diminue ; c'est *l'effet Peltier* qui peut être utilisé pour refroidir des dispositifs électroniques. Si, d'un autre côté, on maintient les deux jonctions à des températures différentes, on observe une différence de potentiel aux bornes du dispositif ; c'est *l'effet Seebeck*.

5.2. Jonction métal-semiconducteur

Lorsqu'un métal est mis en contact avec un semi-conducteur, les niveaux de Fermi s'équilibrent. Deux cas peuvent se présenter selon la valeur relative du travail d'extraction du métal et du semi-conducteur. Pour la discussion, nous allons considérer un semi-conducteur de type n dont le travail d'extraction vaut W et un métal dont le travail d'extraction vaut W_m . Les deux situations à considérer sont celles pour lesquelles on a $W_m > W$ et $W_m < W$. Dans la première, on a un effet redresseur analogue à celui observé dans le cas d'une diode pn ; c'est la diode de Schottky. Dans la seconde ($W_m < W$), on a un contact ohmique.

5.2.1. Diode Schottky

La diode Schottky est constituée d'un métal en contact avec un semi-conducteur. Les deux matériaux sont choisis tels que $W_m > W$. Lorsque les deux éléments sont séparés, on a le diagramme en énergie schématisé dans la figure 5. 2(a). L'affinité électronique du semi-conducteur χ est l'énergie nécessaire pour arracher un électron situé dans le bas de la bande de conduction. Si E_{vide} est l'énergie du vide, le travail d'extraction du métal est défini par $W_m = E_{vide} - E_{(F_m)}$ et celui du semi-conducteur par $W = E_{vide} - E_F$. L'affinité vaut $\chi = E_{vide} - E_c$ où E_c est l'énergie du minimum de la bande de conduction.

Lors de l'évaporation du métal sur le semi-conducteur, les niveaux de Fermi s'alignent et l'on obtient le diagramme montré schématiquement sur la figure 5.2(b). Pour le cas que nous considérons ici, des électrons passent du semi-conducteur vers le métal. Les niveaux d'énergie du semi-conducteur **n** se décalent vers le bas d'une quantité égale à $\Delta W = W_m - W$. C'est précisément la hauteur de la barrière que voit un électron situé dans le bas de la bande de conduction du semi-conducteur. La hauteur de la barrière est différente vue du côté du métal. En effet, il faut rajouter, à ΔW , la quantité

$E_C - E_F$ puisque le niveau de Fermi du métal correspond à la plus grande valeur de l'énergie d'un électron. On a donc, pour le métal, une barrière égale à $\Delta W_m = W_m - \chi$

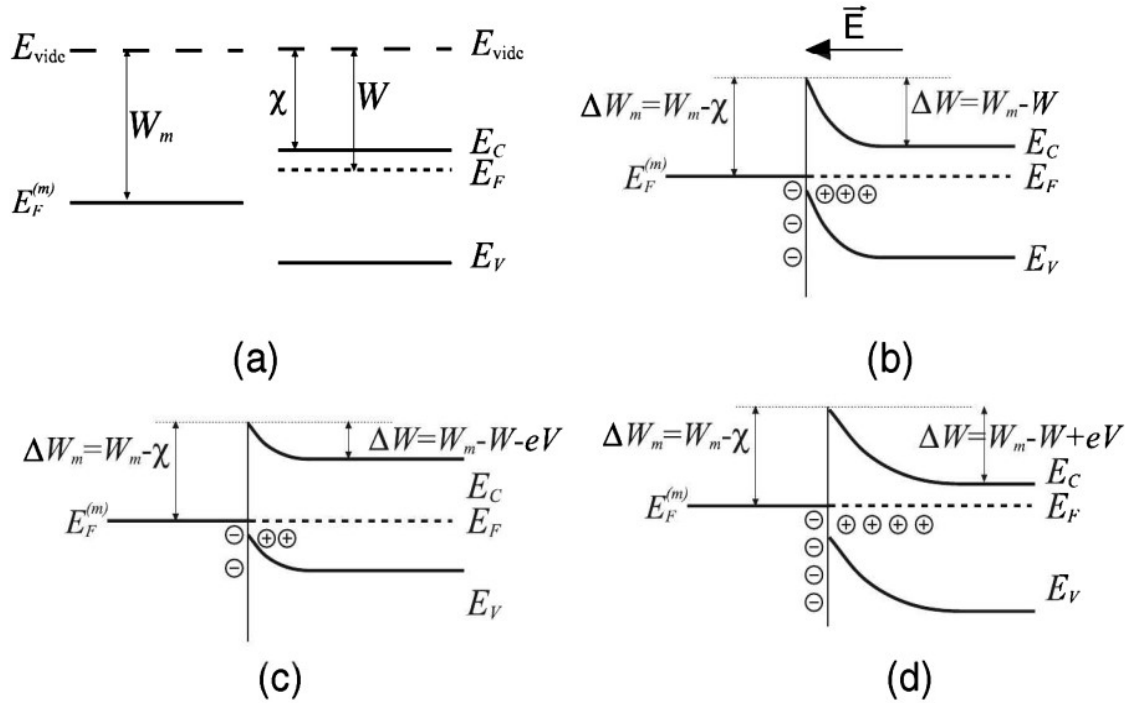


Figure 5.2. Principe de formation d'une jonction métal-semi-conducteur dans le cas où $W_m > W$. La figure (a) montre le diagramme d'énergie des deux matériaux séparés. La figure (b) correspond au cas où la jonction métal-semi-conducteur s'est formée, il s'agit là d'une jonction Schottky. Les figures (c) et (d) correspondent au cas où l'on applique une tension extérieure.

Lors de la formation de la jonction métal-semi-conducteur, des électrons vont donc, au voisinage de la jonction, passer du semi-conducteur vers le métal jusqu'à ce que les niveaux de Fermi soient égaux. Les atomes, chargés positivement, ne migrent bien évidemment pas. On va ainsi créer, au voisinage de la jonction, une zone de déplétion dans le semi-conducteur, appauvrie en porteurs de charges libres. Par contre, au niveau du métal, il se produit une accumulation de charges négatives (les électrons) en surface ; il n'y a pas de zone de déplétion. À l'équilibre, *i.e.* lorsque les potentiels chimiques sont égaux, il y a, au niveau de la jonction, un champ électrique dirigé du semiconducteur vers le métal, qui s'oppose à un déplacement supplémentaire d'électrons. Toutefois, à cause de l'agitation thermique, il y a toujours un flux d'électrons passant du semi-conducteur vers le métal et vice versa mais ces flux sont égaux à l'équilibre ce qui nous fait dire, au point de vue macroscopique, qu'il n'y a pas de transfert d'électrons.

Si nous appliquons maintenant une tension positive, V , au métal par rapport au semi-conducteur, nous diminuons la hauteur de la barrière s'opposant au passage des électrons du semi-conducteur vers le métal de eV (figure 5.2 (c)). Elle devient $\Delta W = W_m - W - eV$.

Par contre, la barrière pour les électrons qui passent du métal vers le semi-conducteur reste inchangée car il n'y a pas de chute de potentiel dans un métal. Par conséquent, le flux d'électrons passant du semi-conducteur vers le métal augmente alors que le flux inverse reste inchangé. Ce flux d'électrons du semi-conducteur vers le métal, qui se traduit par un courant allant du métal au semiconducteur, augmente exponentiellement avec la tension appliquée. La jonction est dite polarisée dans le sens « passant ».

Si nous appliquons une tension positive, V , au semi-conducteur par rapport au métal, la barrière s'opposant aux électrons est augmentée de eV et devient $\Delta W = W_m - W + eV$ (figure 5.2 (d)).

La zone de déplétion s'élargit et le flux d'électrons passant du semi-conducteur au métal diminue fortement. Par contre, le flux d'électrons du métal vers le semi-conducteur reste inchangé. Le passage d'électrons du semi-conducteur vers le métal est pratiquement bloqué et l'on n'observe qu'un très faible courant qui sature à une valeur pratiquement indépendante de V . La jonction est dite polarisée dans le sens « bloquant ».

- Le courant inverse est dû aux porteurs minoritaires qui diffusent à travers la zone de déplétion dans le cas d'une jonction ***p-n*** alors qu'il provient des porteurs majoritaires qui franchissent la barrière, en provenance du métal, dans le cas d'une jonction Schottky. Par conséquent, il en résulte une forte dépendance en température pour la jonction *pn*.
- Le courant dans le sens passant provient de l'injection de porteurs minoritaires venant des zones *p* et *n* pour la jonction *pn* alors qu'il provient des porteurs majoritaires en provenance du semiconducteur pour la diode Schottky.
- La tension d'opération dans le sens passant est inférieure pour une diode Schottky à celle d'une jonction *pn*.
- La diode Schottky commute beaucoup plus rapidement que la jonction *pn*.

5.2.2. Jonction Ohmique

Considérons un semi-conducteur de type *n* et plaçons-nous maintenant dans le cas où $W_m < W$, *i.e.* où le travail d'extraction est inférieur à celui du semiconducteur. Au contact, les niveaux de Fermi s'égalisent et il n'y a pas de barrière empêchant les électrons de passer. Ceci permet, en principe,

d'avoir un courant dans les deux directions sans avoir de perte appréciable. On qualifie pour cette raison ce type de contact de *jonction ohmique*.

Les jonctions ou contacts ohmiques sont importants dans la technologie des semiconducteur, notamment pour les interconnexions. Bien qu'en principe il suffise, comme nous venons de le voir, d'avoir un métal dont le travail d'extraction est inférieur à celui du semi-conducteur, la réalisation pratique de contacts purement ohmiques est difficile et demande un bon savoir-faire. Dans la technologie des semi-conducteurs (micro-électronique), le métal utilisé pour faire les interconnexions est l'aluminium. Toutefois, depuis peu de temps, on sait également mettre en œuvre du cuivre. Pour réaliser un bon contact ohmique, on crée une fine zone n^+ , *i.e.* fortement dopée n , au voisinage du contact métallique. Cette couche de très faible épaisseur réduit la largeur de la zone de déplétion et permet aux électrons de passer par effet tunnel.

5.3. Transistors bipolaires

5.3.1. Introduction

Les transistors sont les composants de base de l'électronique. Il s'agit de tripôles, permettant de faire passer un courant entre deux de ses bornes que l'on contrôle par une tension ou un courant appliqué à la troisième borne.

- ◆ Les transistors bipolaires, constitués de trois zones de semi-conducteurs de type NPN ou de type PNP.
- ◆ Les transistors unipolaires, basés sur un seul type de porteurs de charge dans le passage du courant.

5.3.2. Définitions

Un transistor bipolaire est constitué d'un monocristal de semi-conducteur (principalement le silicium), dopé pour obtenir deux jonctions, disposées en série et de sens opposé. Il existe donc deux types fondamentaux de transistors bipolaires : les transistors NPN dans lesquels une mince couche de type P est comprise entre deux zones de type N, figure 5.3.a, et les transistors PNP dans lesquels une mince couche de type N est comprise entre deux zones de type P, figure 5.3.b. Les différentes parties de chaque transistor sont :

- **L'émetteur (E)** est la couche la plus fortement dopée. Son rôle est d'injecter des électrons dans la base.

- **La base (B)** est la couche la plus faiblement dopée et très mince. Elle transmet au collecteur la plupart des électrons venant de l'émetteur.
- **Le collecteur (C)** recueille les électrons qui lui viennent de la base d'où son nom.

Dans cette partie, l'étude sera menée sur le transistor bipolaire NPN parce qu'il est le plus utilisé et le plus facile à réaliser. Le fonctionnement de l'autre type de transistor se déduit en échangeant les rôles des électrons ainsi que des trous et en inversant les signes des tensions d'alimentation et des courants.

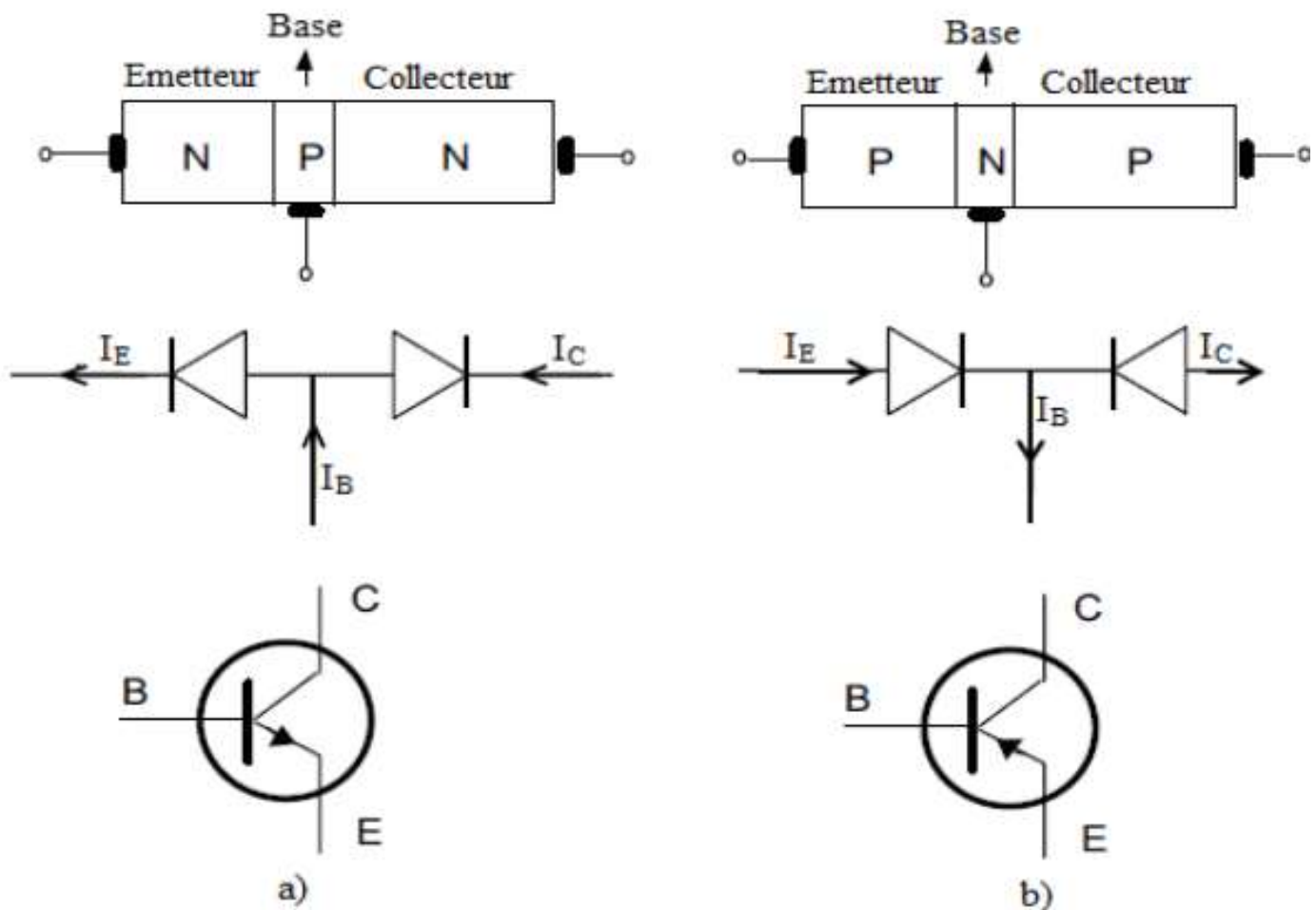


Figure 5.3 : Représentations schématiques et symboles des transistors bipolaires, a) transistor bipolaire NPN, b) transistor bipolaire PNP.

5.3.1. Transistor NPN non polarisé

Si les régions NPN du transistor ne sont plus isolées l'une de l'autre, les électrons libres diffusent à travers les deux jonctions ce qui donne deux zones de déplétion (figure 5.4). Ces deux zones de transition sont dépourvues de porteurs majoritaires et la barrière de potentiel pour chacune d'elles est d'environ 0,6 à 0,7 Volt (cas du silicium). Puisque les trois régions ne sont pas dopées de la même manière, les deux zones de déplétion auront donc des épaisseurs différentes. Ainsi, la zone de déplétion pénètre peu dans l'émetteur qui est fortement dopé mais profondément dans la base qui est très peu dopée. Du côté du collecteur, la pénétration de la zone de déplétion sera moyenne.

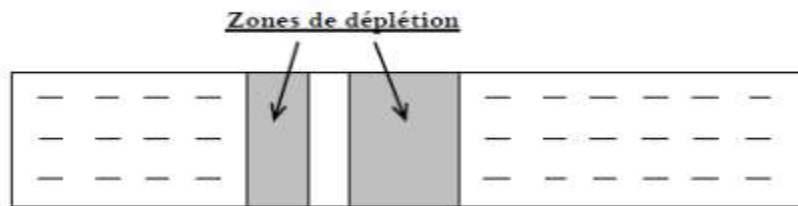


Figure 5.4 : Transistor bipolaire NPN non polarisé.

5.3.2. Transistor NPN polarisé

Parmi les différentes façons de polariser un transistor de type NPN, une seulement, présente un intérêt primordial. C'est la polarisation de la jonction **émetteur-base** en direct et la jonction **collecteur-base** en inverse. Dans ce cas, nous obtenons la configuration de la figure 5.5.

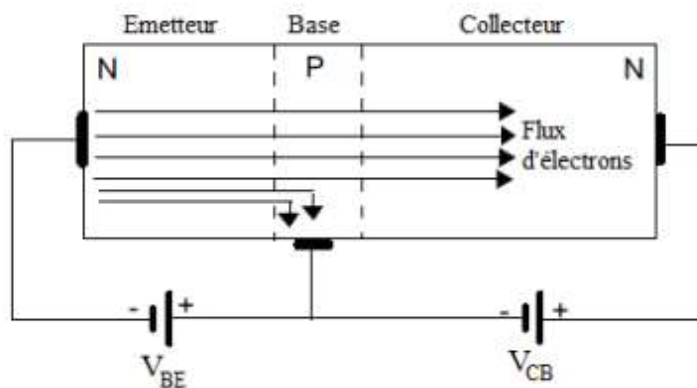


Figure 5.5 : Polarisation d'un transistor NPN et principe de l'effet transistor.

5.3.2.1. Courants à travers les jonctions

On mesure les courants entre deux électrodes reliées à un générateur quand la troisième est déconnectée.

- a) **Jonction Base-Emetteur** : Si on polarise la jonction B-E en direct, le courant I_{BE} est intense. Par contre en polarisation inverse, I_{EB0} est très faible (courant de saturation). Voir figure 5.6.

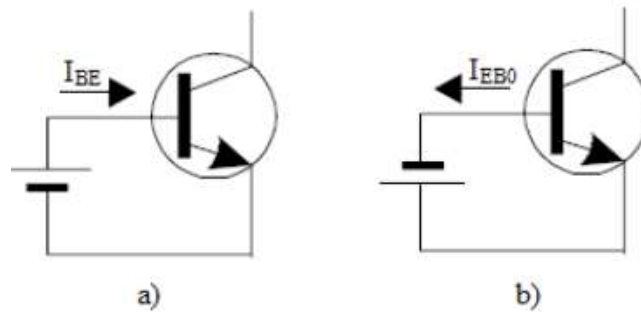


Figure 5.6 : Jonction Base-Emetteur polarisée, a) en direct, b) en inverse.

- b) **Jonction Base-Collecteur** : Si on polarise la jonction B-C en direct, le courant I_{BC} est intense, voir figure 6.7.a. Par contre en polarisation inverse, I_{CB0} est très faible (courant de saturation). En effet, le dopage du collecteur étant faible celui-ci contient peu de porteurs libres.

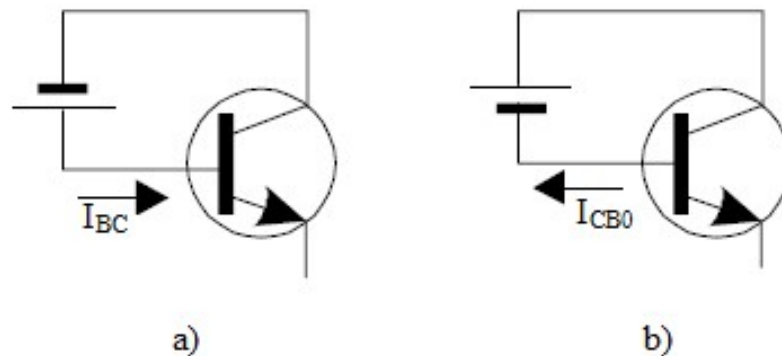


Figure 5.7 : Jonction Base-Collecteur polarisée, a) en direct, b) en inverse.

- c) **Espace Emetteur-Collecteur** :

Si la jonction B-E est polarisée en inverse, I_{EC0} est très faible, mais on a : $I_{EC0} > I_{EB0}$. Si la jonction B-E est polarisée en direct, on mesure un courant I_{CE0} très faible avec : $I_{CE0} > I_{EB0} \gg I_{CB0}$.

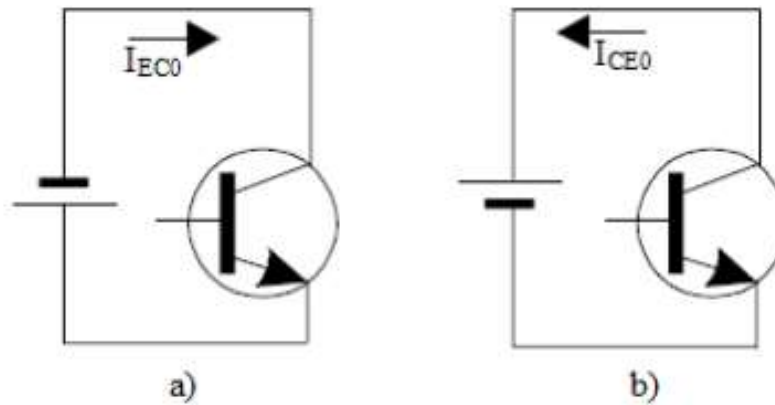


Figure 5.8 : a) Jonction Base-Emetteur polarisée en inverse, b) Jonction Base-Emetteur polarisée en direct.

5.3.2.2. Effet transistor : Gains en courant

L'effet transistor consiste à moduler le courant inverse de la jonction B-C, polarisée en inverse, par une injection de porteurs minoritaires (électrons) dans la base, à partir de la jonction E-B polarisée en direct. Les porteurs minoritaires, injectés dans la base, sont ensuite soumis à un champ intense de la jonction B-C, polarisée en inverse, puis dérivent vers le collecteur. Pour que ces porteurs atteignent la jonction B-C, l'épaisseur de la base doit être inférieure à leur longueur de diffusion. Cette condition est fondamentale pour éviter la recombinaison des porteurs minoritaires lors de la traversée de la base.

En premier lieu, supposons que seule la jonction B-C soit polarisée en inverse ($V_{CB} \neq 0$) et la jonction E-B est en court-circuit ($V_{BE} = 0$), figure 5.9. Dans ce cas, la jonction B-C est traversée par un courant très faible dû aux porteurs minoritaires (électrons), appelé I_{CB0} .

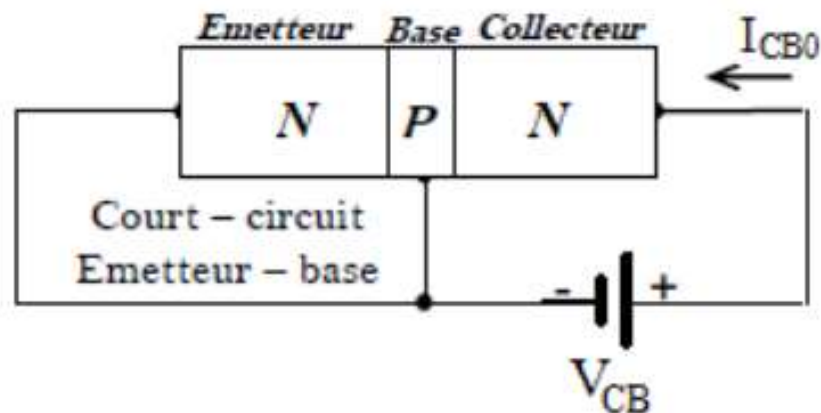


Figure 5.9: Montage en court-circuit de la jonction E-B.

Polarisons maintenant la jonction B-E en direct, figure 5.10. Les électrons qui sont majoritaires dans la région de l'émetteur (type N) diffusent en grande quantité à travers la jonction E-B, polarisée en direct, créant ainsi un courant émetteur I_E . Les électrons de l'émetteur traversent en majorité la base et arrivent jusqu'au collecteur. Ainsi l'émetteur « injecte » ou « émet » des porteurs majoritaires et le collecteur les collecte.

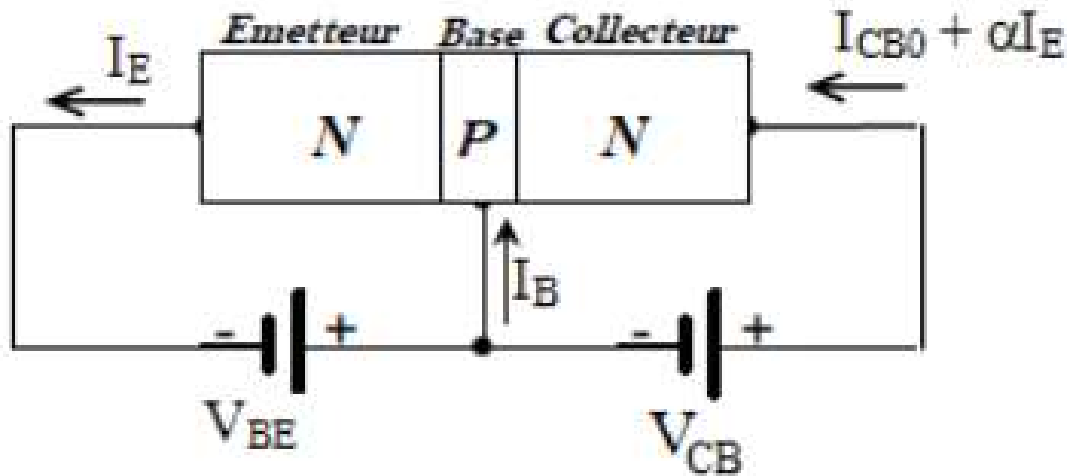


Figure 5.10 : Montage montrant la polarisation des jonctions B-E en direct et C-B en inverse.

Nous appelons α la proportion des électrons dans le cas du transistor NPN émis par l'émetteur, qui parviennent jusqu'au collecteur ; α est généralement proche de l'unité. Le courant total devient :

$$I_C = \alpha I_E + I_{CB0}$$

Et puisque la loi des nœuds de Kirchhoff nous permet d'écrire :

$$I_E = I_C + I_B$$

Alors, en éliminant I_E on obtient :

$$I_C = \frac{\alpha}{1-\alpha} I_B + \frac{1}{1-\alpha} I_{CB0}$$

$$\beta = \frac{\alpha}{1-\alpha}$$

On obtient :

$$I_C = \beta I_B + (\beta + 1)I_{CB0} \approx \beta I_B$$

Cette dernière relation caractérise l'effet transistor : en injectant un courant I_B très faible dans la base, nous commandons un courant de collecteur I_C beaucoup plus intense

5.3.3. Réseau des caractéristiques statiques du transistor NPN

Le transistor comporte trois accès et caractérisé par six grandeurs électriques :

- 3 courants I_C , I_B et I_E

- 3 tensions V_{CE} , V_{BE} , V_{BC}

Mais $I_E = I_C + I_B$ et $V_{CB} = V_{CE} - V_{BE}$, donc quatre relations indépendantes sont nécessaires pour le caractériser.

Généralement, le transistor qui est un composant à trois bornes est utilisé en tant que quadripôle amplificateur. Dans la plupart des applications, une des bornes est commune à l'entrée et à la sortie.

On a donc trois possibilités de montage de transistor :

- Montage base commune utilisé en haute fréquence.
- Montage collecteur commun utilisé en adaptation d'impédance.
- Montage émetteur commun utilisé en amplification et le plus commun.

5.3.3.1. Montage émetteur commun

Les bornes d'entrée du quadripôle sont la base et l'émetteur ; les grandeurs d'entrée sont : I_B et V_{BE} .

La sortie se fait entre le collecteur et l'émetteur ; les grandeurs de sortie sont : I_C et V_{CE} , figure 5.11.

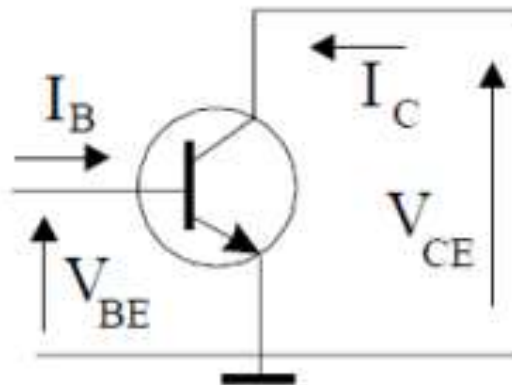


Figure 5.11 : Montage émetteur commun.

- ✓ **Caractéristique d'entrée** : Pour décrire la caractéristique d'entrée, on mesure le courant d'entrée I_B en fonction de la tension d'entrée V_{BE} pour diverses tensions de sorties V_{CE} . Cette caractéristique a l'allure de la courbe $I=f(V)$ de la jonction polarisée en direct, de forme exponentielle :

$$I_B = I_S \left(e^{\frac{qV_{BE}}{k_B T}} - 1 \right)$$

Le réseau d'entrée $I_B = f(V_{BE})$ se réduit dans la pratique à une seule courbe, toutes les courbes étant confondues.

- ✓ **Caractéristique de transfert en courant** : C'est la caractéristique $I_C = f(I_B)$ pour des valeurs constantes de V_{CE} . On a vu que le courant collecteur s'écrit :

✓

$$I_C = \frac{\alpha}{1-\alpha} I_B + \frac{1}{1-\alpha} I_{CB0}$$

$$I_{CE0} = \frac{1}{1-\alpha} I_{CB0}$$

Le courant I_{CE0} peut être mesuré en déconnectant la base ($I_B = 0$) avec une tension V_{CE} positive de sorte que les jonctions soient normalement polarisées. C'est un courant de l'ordre d'une centaine de μA . Ce courant est donc généralement négligeable devant I_C qui est de l'ordre de la mA.

La caractéristique $I_C = f(I_B)$ ne passe pas par l'origine, mais par I_{CE0} , qui est très proche de l'origine. Elle présente une légère courbure pour des courants du collecteur très faibles et tend vers une droite dès que le courant I_C dépasse quelques centaines de μA .

- ✓ **Caractéristique de sortie** : C'est la caractéristique $I_C = f(V_{CE})$ à I_B constant. Dès que la tension V_{CE} devient suffisamment positive pour que les deux jonctions soient normalement polarisées, on a :

$$I_C = \beta I_B + I_{CE0}$$

Si le gain β en courant était rigoureusement constant, les caractéristiques seraient des droites horizontales et équidistantes pour des intervalles ΔI_B égaux.

Les caractéristiques $I_C = f(V_{CE})$ sont limitées par deux zones proches des axes qu'il n'est pas possible d'utiliser :

- Une zone de saturation qui est proche de l'axe des courants I_C . Cette zone correspond au cas où les deux jonctions sont polarisées en direct ;
- Une zone de blocage qui est proche de l'axe des tensions V_{CE} et obtenue pour un courant de base nul ($I_B = 0$). Il n'est pas possible d'avoir I_B négatif à cause de la jonction B-E.

✓ **Caractéristique de transfert en tension :**

C'est la caractéristique $V_{BE} = f(V_{CE})$ à I_B constant. Cette caractéristique présente peu d'intérêt et se présente sous forme de droites pratiquement horizontales. La tension B-E ne dépend pratiquement pas de la tension entre le collecteur et l'émetteur dès que cette tension dépasse 1 volt.

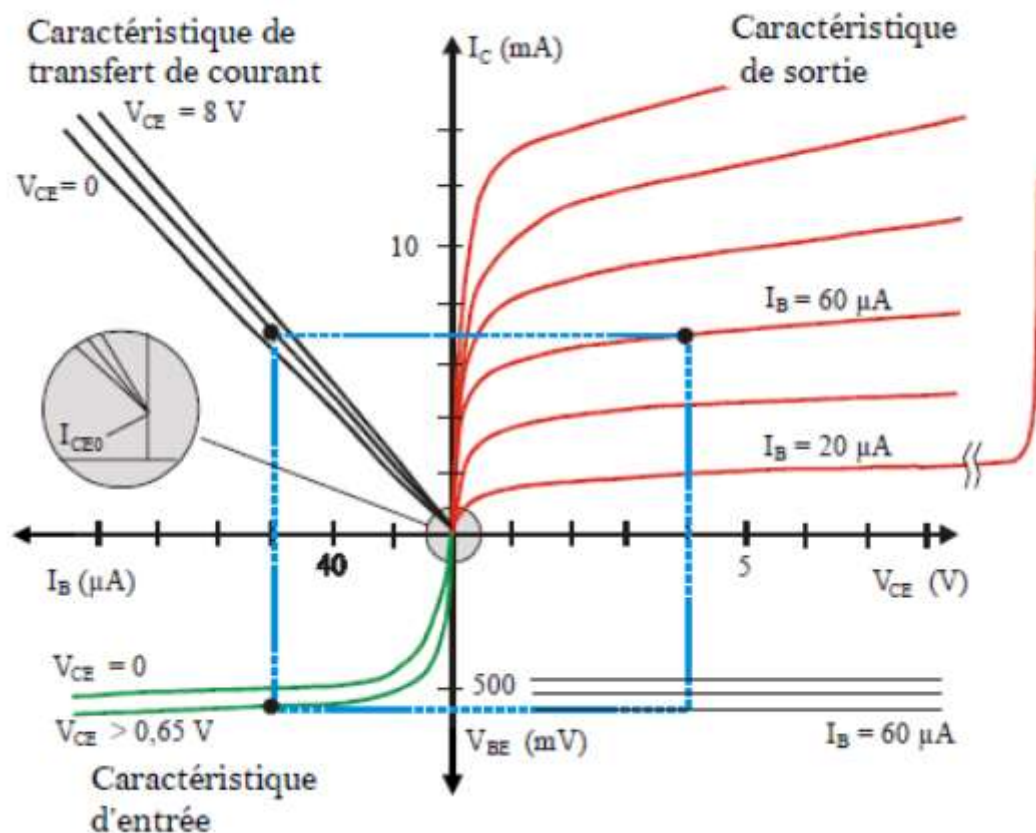


Figure 5.12 : Réseau de caractéristiques statiques du transistor NPN.

5.3.4. Régimes de fonctionnement du transistor

Nous allons chercher à déterminer la valeur du courant I_C et de la tension V_{CE} en fonction des éléments du montage de la figure 5.13. Nous disposons de deux équations :

- L'une provenant du transistor, donnée par le réseau de caractéristiques $I_C = f(V_{CE})$ à I_B constant,
- L'autre de la loi des mailles, soit : $V_{CC} = R_C I_C + V_{CE}$

La droite représentative de cette équation est appelée droite de charge statique. Traçons cette droite dans le système d'axes $I_C = f(V_{CE})$, figure 5.14.

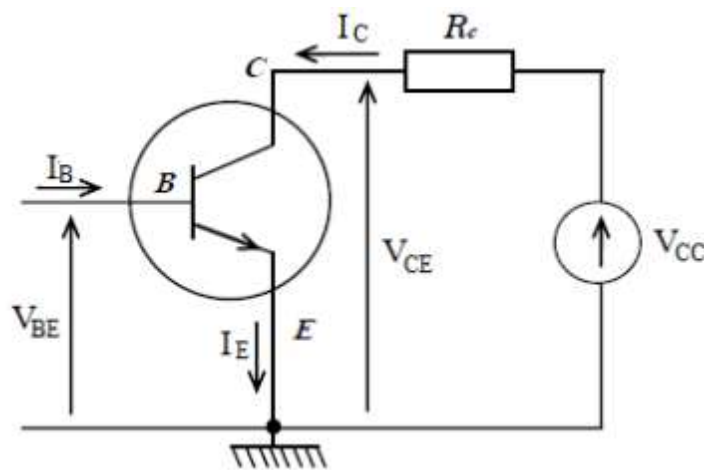


Figure 5.13 : Montage en émetteur commun.

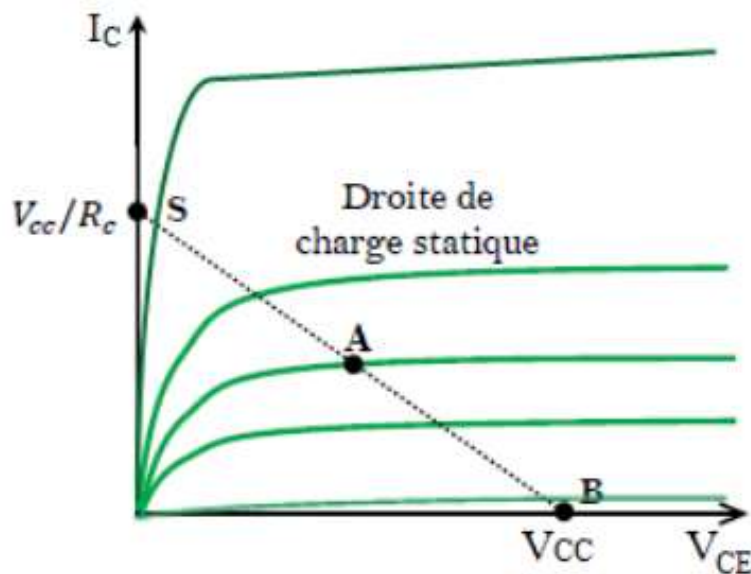


Figure 5.14 : Représentation graphique du point de fonctionnement : droite de charge statique.

Les valeurs de I_C et de V_{CE} sont les coordonnées du point d'intersection de la droite de charge statique et de la caractéristique $I_C = f(V_{CE})$ correspondant à la valeur de I_B imposée par le réseau d'entrée.

Nous distinguons trois positions remarquables correspondant à trois modes de fonctionnement particulier du transistor :

- Le point A dans la partie linéaire et horizontale des caractéristiques. Ce point correspond à un fonctionnement linéaire en amplification. Le courant I_C se déduit de la droite de charge statique par la relation $I_C = \beta I_B$.
- Le point S dans la partie montante des caractéristiques. Le transistor est saturé, $V_{CE} \approx 0$. Toute augmentation de I_B est pratiquement sans effet sur la valeur de I_C . Le transistor se comporte, entre collecteur et émetteur, comme un interrupteur fermé. On note $V_{CE} = V_{CEsat}$.
- Le point B pratiquement sur l'axe des V_{CE} (I_C est très faibles). Le transistor est bloqué. Il se comporte, entre émetteur et collecteur, comme un interrupteur ouvert.

5.3.5. Limitations physiques

Comme tout composant électronique, le transistor bipolaire ne peut fonctionner que dans une zone bien déterminée. Un certain nombre de limitations physiques doivent être prises en compte pour assurer un fonctionnement sûr du composant.

5.3.5.1. Tensions de claquage

Nous avons vu l'existence des courants de fuites I_{CE0} (en base commune et émetteur ouvert) et I_{CB0} (en émetteur commun et base ouverte). Si on augmente exagérément les tensions, les courants de fuites augmentent par effet avalanche et peuvent être la cause de la destruction de transistor par échauffement. Ces tensions à ne pas dépasser sont données par le constructeur et sont généralement notées BV_{CE0} et BV_{CB0} (BV est l'abréviation de Breakdown Voltage).

5.3.5.2. Limitation en puissance

La puissance dissipée par un transistor au repos est donnée par la formule suivante :

$$P = V_{BE} I_B + V_{CE} I_C \approx V_{CE} I_C < P_{max}$$

Cette puissance est limitée à cause de l'échauffement du transistor. La température maximale de la jonction ne doit pas dépasser 200 °C dans le cas du silicium. Le lieu limite est donc une hyperbole dans le plan de sortie I_C, V_{CE} , figure 5.14.

5.3.5.3. Limitation du courant maximum

Le courant maximum du collecteur doit rester inférieure à une certaine valeur I_{Cmax} sous peine de destruction du transistor.

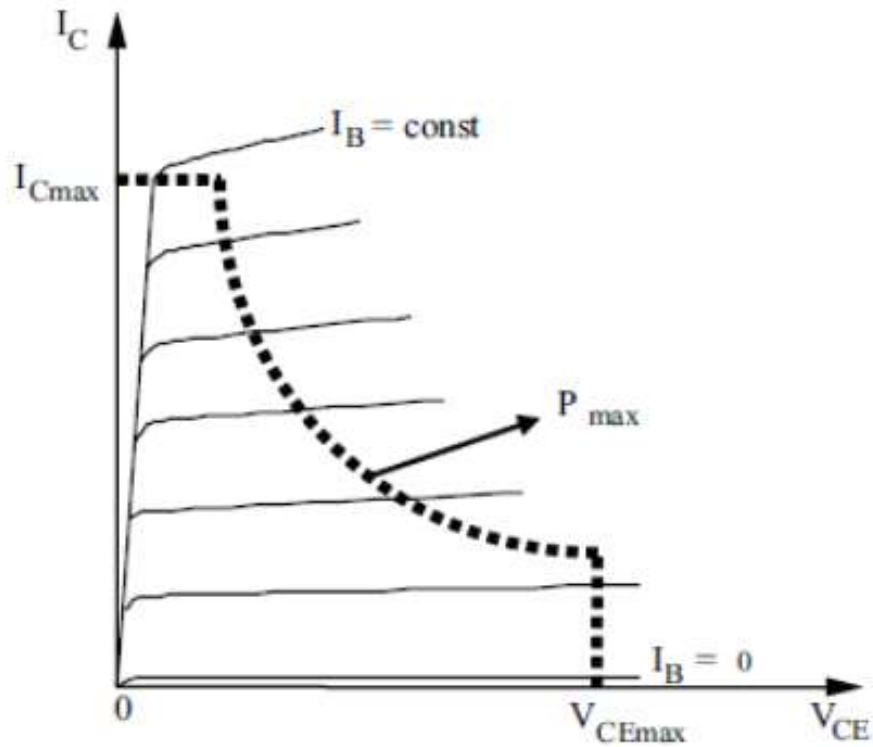


Figure 5.15 : Zone de fonctionnement d'un transistor.

Références bibliographiques

1. Henry Mathieu, Hervé Fanet, Cours et exercices corrigés « Physique des semi-conducteurs et des composants électroniques », 6^{ème} édition, Dunod, Paris, 2009.
2. Peter Y. Yu Manuel Cardona, Fundamentals of Semiconductors, 3rd Edition, 2nd Corrected Printing Springer Berlin Heidelberg New York, 2005.
3. Christian Ngô, Hélène Ngô, Cours et exercices corrigés « Physique des semi-conducteurs » 4^{ème} édition, Dunod, Paris ,2012.
4. Kittel C., Physique de l'état solide, 7^{ème} Edition. Dunod, Paris,2003.
5. Walter A. Harrison, Solid State Theory, Dover publications, Inc, New york, 1980.
6. Walter.A. Harrison, Electronic structure and the properties of solids, Ed. W.H. Freeman and company - San Francisco (1980).
7. Neil W. Ashcroft, N. David Mermin, Physique des solides, EDP Sciences, 2002.