

République algérienne Démocratique et Populaire
Université de A/MIRA de Bejaïa
Faculté des Sciences exactes
Département de physique

Mémoire de fin d'études

En vue de l'obtention du diplôme de Master en Physique théorique

Promoteur : ZENIA Hand

Étude des propriétés d'un électron dans un réseau cristallin à une dimension

IKHENACHE Nadira

Année universitaire: 2011 - 2012

Remerciements

Je remercie mon promoteur M. Zenia Hand.

Je tiens à exprimer ma profonde gratitude à M. Adel Kassa pour l'excellente formation dont nous avons bénéficié que ça soit en physique ou en informatique, je tiens aussi à le remercier pour tout le temps qu'il m'a accordé en détriment de ses occupations.

Je conclurai en remerciant chaleureusement ma formidable famille en particulier les deux êtres les plus chères à mon cœur : Mes parents !

Sommaire

Introduction	1
I Approche théorique	4
0.1 Théorie des bandes	5
0.1.1 Généralités	5
0.1.2 Le concept de pseudo-potentiel	6
0.1.3 Théorème de Bloch	7
0.1.4 Modèle de Krönig-Penney	10
0.1.5 Conducteurs, semiconducteurs et isolants	14
0.2 Réseaux irréguliers : La localisation d'Anderson	18
II Approche numérique	23
0.3 Résolution de l'équation de Schrödinger à une dimension pour un électron dans un réseau périodique de puits carrés de potentiel	24
0.3.1 Méthode directe	27
0.3.2 Méthode de Numerov(differences finies améliorées)	28
0.3.3 Résultats	30
0.4 Introduction d'une impureté dans le réseau régulier	36
0.5 Réseau désordonné	39
Conclusion	45
Appendices	46
.1 La méthode de Dichotomie ou Bissection	47
.2 La méthode de Numerov	48

Table des figures

1	Le potentiel périodique ressenti par un électron dans un réseau unidimensionnel d'atomes (ligne mauve) et la schématisation avec des puits de potentiel carrés (ligne bleue).	11
2	Les niveaux d'énergie pour un réseau régulier de dix puits de potentiel (obtenus avec la méthode directe)	31
3	Les fonctions d'ondes des deux derniers états de la dernière bande	32
4	Les fonctions d'ondes et les probabilités de présence de l'électron correspondant à un niveau d'énergie de chaque bande, obtenues avec la méthode directe	33
5	Les fonctions d'ondes et les probabilités de présence de l'électron correspondants à l'état fondamental et les trois premiers états excités des méthodes directe et Numerov	35
6	Les niveaux d'énergie du réseau avec impuretés	37
7	Les fonctions d'ondes des niveaux situés à l'intérieur des bandes	38
8	Les fonctions d'ondes et probabilités de présence des états situés dans les gaps	39
9	Les niveaux d'énergie pour un réseau irrégulier de dix puits de potentiel .	40
10	Les fonctions d'ondes et les probabilités de présence de l'électron correspondants aux premiers niveaux des trois premières bandes d'énergie . . .	42
11	Les fonctions d'ondes et les probabilités de présence de l'électron correspondant aux niveaux d'énergie des dernières bandes	43
12	La localisation des états des bords des bandes	44

Résumé

Dans ce travail nous avons tenté d'étudier les propriétés d'un électron dans un cristal à une dimension, et cela en utilisant deux approches. Une approche théorique, dans laquelle nous avons mis le point sur la théorie des bandes. Ainsi, nous avons donné un aperçu sur la formation des bandes d'énergie dans un solide. Puis, nous avons pris le cas d'un potentiel périodique de type Kronig-Penney. Ensuite, nous avons essayé d'expliquer comment la structure de bande d'un solide peut déterminer sa nature. Nous avons aussi abordé le concept de pseudo-potentiel ainsi que celui de la localisation d'Anderson.

La deuxième approche est numérique. Ici nous avons vérifié les résultats prédis par la théorie en résolvant l'équation de Schrödinger pour un électron se mouvant dans un potentiel périodique.

Mots clés : Théorie des bandes, Kronig-Penney, Théorème de Bloch, réseau périodique, localisation, pseudo-potentiel.

Introduction

Les tentatives de résolution de l'équation de Schrödinger pour un réseau régulier de puits de potentiel datent des premiers jours de la mécanique quantique. En 1928, Bloch[1] a montré que les fonctions propres d'un réseau périodique infini sont de la forme d'ondes planes modulées par une fonction qui a la même périodicité que le réseau. Il a suggéré dans ses travaux qu'il y avait des niveaux d'énergie pour lesquels aucun état physique existe (les gap d'énergie) et des bandes quasi-continues de niveaux d'énergie. La même conclusion a été atteinte de manière indépendante en utilisant l'approximation des liaisons fortes, dans laquelle les niveaux d'énergie discrets d'un électron dans un seul puits sont "éclatés" en bandes quasicontinues, et cela par la présence de puits voisins .

En 1930, Kronig et Penney[2] ont publié leur solution pour un réseau infini composé de puits de potentiel carrés équidistants. Un élément clé de leur analyse était que les solutions devaient être des ondes de Bloch, une hypothèse valable uniquement pour un réseau strictement infini. Ils ont donc invoqué les conditions aux limites périodiques, plutôt que d'exiger que la fonction d'onde doit être nulle (plutôt tendre vers zéro) aux extrémités du réseau. Leur résultat a donné une relation entre le vecteur d'onde de Bloch et l'énergie de l'état sous la forme d'une équation transcendante impliquant les paramètres du réseau. Encore une fois, ils ont prédit des bandes d'énergie et les écarts entre elles, et ils ont fourni un moyen général pour calculer les limites entre les deux.

En 1950, avec l'apparition et l'usage intensif des semiconducteurs, l'attention était dirigée vers les réseaux irréguliers. On a constaté que le fait de changer le potentiel dans certains sites du réseau (équivalent à un cristal légèrement dopé) conduisait souvent à des niveaux d'énergie qui se trouvaient au milieu des gaps précédemment inoccupés. On a soupçonné que ces nouveaux états étaient fortement localisés sur les irrégularités qui les génèrent (par opposition aux ondes de Bloch qui s'étendent à travers le cristal entier), cette propriété est un élément clé dans l'étude de la conductivité des solides. Ce ne fut cependant pas prouvé jusqu'à cet article précurseur écrit par Anderson[3] en 1958. Un peu plus tard, quand les solides amorphes sont devenus importants, les réseaux avec un faible degré de désordre étaient largement étudiés.

Dans la partie théorique de ce mémoire, nous avons essayé de donner les concepts de

base du comportement d'un électron dans un solide. Ainsi, nous avons tenté de donner un aperçu sur la théorie des bandes d'énergie dans les solides cristallins. L'importance de cette théorie réside dans le fait que la majorité des propriétés physiques des matériaux peuvent être expliquées en utilisant leurs structures de bande.

En général, la structure de bande d'un solide peut être obtenue en résolvant l'équation de Schrödinger pour un solide qui contient un grand nombre d'atomes et d'électrons en interaction. Pour simplifier la tâche de résolution d'un problème à N-corps, nous sommes amenés à effectuer plusieurs approximations. En premier lieu, on commence par négliger le mouvement des noyaux, en supposant qu'ils sont figés sur leur position d'équilibre sur chaque site du réseau, "c'est l'approximation de Born Oppenheimer¹". De cette façon, leurs coordonnées apparaissent dans le problème comme de simples paramètres. Cependant, le problème demeure un problème à plusieurs électrons qu'on ne peut pas résoudre explicitement. C'est pour cette raison qu'on a besoin d'utiliser une approximation supplémentaire pour résoudre l'équation de Schrödinger. L'une des méthodes les plus efficace pour résoudre un problème à plusieurs électrons dans un cristal est l'approximation à un électron. Dans cette méthode la fonction d'onde totale des électrons est choisie comme une combinaison linéaire des fonctions d'ondes individuelles, dans laquelle chaque fonction contient uniquement les coordonnées d'un seul électron, en considérant ces derniers indépendants les uns des autres. C'est cette approximation qui sert de cadre de base pour calculer la structure de bande dans les solides.

Ensuite, un autre concept est abordé, celui de pseudo-potentiel[4]. En général, calculer de manière exacte la structure de bande d'un solide est très complexe. Ce qui pose problème est que le potentiel du solide est assez compliqué, si on prend en considération l'effet induit par chaque particule. La méthode du pseudo-potentiel vient remplacer les effets du noyau et des électrons du cœur par un potentiel effectif qui interagit uniquement avec les électrons de valence, et cela simplifie considérablement le problème.

Plus loin dans ce travail, nous allons traiter le cas particulier d'un potentiel périodique. Nous allons voir que la fonction d'onde doit obéir au théorème de Bloch[5]. Ensuite, nous nous essayerons à la résolution de l'équation de Schrödinger pour une particule se mouvant dans un potentiel en créneau de type Kronig-Penney[6][4]. Puis, nous allons expliquer pourquoi la structure de bande d'un matériau détermine sa nature, c'est-à-dire : isolant, semiconducteur ou métal. Enfin, nous allons aborder très brièvement un phénomène très important dans le domaine de la physique de la matière condensée, celui de la localisation d'Anderson.

Dans la partie numérique, nous avons tenté de vérifier si les résultats théoriques sont

1. L'approximation de Born Oppenheimer est uniquement valide lorsque l'état fondamental électronique est non dégénéré

maintenus, en élaborant des programmes simulant le mouvement d'un électron dans un réseau de puits carrés de potentiel. Ainsi, nous allons voir si la structure de bande est réellement obtenue. Puis, nous allons modifier le réseau pour simuler l'effet d'une impureté dans un cristal pour voir les modifications apportées par un dopage. Finalement, nous allons essayer d'observer le phénomène de localisation, en rendant le réseau complètement désordonné.

Première partie

Approche théorique

0.1 Théorie des bandes

0.1.1 Généralités

Dans un atome isolé, l'énergie des électrons ne peut posséder que des valeurs discrètes et bien définies. Par contraste, dans le cas d'un électron parfaitement libre, elle peut prendre n'importe quelle valeur positive. Dans un solide, la situation est intermédiaire : l'énergie d'un électron peut avoir n'importe quelle valeur à l'intérieur de certains intervalles.

Pour bien comprendre ces propriétés, il est essentiel de commencer par la théorie des orbitales moléculaires. Dans cette théorie de base, on suppose que si plusieurs atomes sont réunis pour former une molécule, les orbitales atomiques² de chaque atome ne vont plus rester indépendantes les uns des autres, mais elles vont se combiner pour former des orbitales moléculaires du même nombre que les orbitales atomiques.

Prenons pour exemple une molécule constituée de deux atomes et chacun des atomes possède une orbitale atomique. Le résultat est qu'il va y avoir deux orbitales moléculaires bien distinctes qui vont apparaître : une orbitale avec une énergie plus basse que celle de l'orbitale atomique de l'atome isolé, et qu'on va appeler orbitale liante, l'autre orbitale moléculaire aura une énergie supérieure à celle de l'orbitale atomique de l'atome isolé et on va l'appeler orbitale anti-liante et les deux sont séparées par un gap d'énergie.

Le concept d'orbitales liantes et anti-liantes peut être étendu au cas d'un grand nombre d'atomes en supposant que les orbitales de chaque atome interagissent avec celles des plus proches voisins. On arrive ainsi à la formation de séries de niveaux discrets très proches les uns des autres que l'on peut assimiler à des bandes. Les orbitales liantes vont former la bande de valence et les orbitales anti-liantes vont former la bande de conduction. Bande de valence et bande de conduction sont séparées par une bande interdite (la bande interdite ne contient aucun état physique)

Dans une telle disposition, les électrons de valence ne sont plus liés à un atome particulier mais bien à l'ensemble du réseau d'atomes. D'une manière générale, dans un solide quelconque, c'est la disposition et le remplissage des bandes permises qui déterminent les propriétés électriques des matériaux, et qui permettent de les classer en conducteurs, isolants et semiconducteurs.

2. Après l'arrivée de la mécanique quantique et la notion de probabilité, les électrons ne sont plus considérés comme gravitant autour du noyau dans une orbite circulaire ou même elliptique, mais occupent de manière probabiliste certaines régions de l'espace autour du noyau et c'est ces régions là qu'on appelle orbitales atomiques.

0.1.2 Le concept de pseudo-potentiel

En physique du solide, le concept de pseudo-potentiel est utilisé comme une approximation pour décrire d'une manière simplifiée un système complexe. Cette méthode a pour principe de substituer le potentiel du noyau et les effets des électrons du cœur par un potentiel effectif plus faible que le potentiel original du réseau, qui agit seulement sur les électrons de valence.

Cette technique vient corriger les failles de l'approximation «cœur gelé »(frozen-core approximation) qui consiste à calculer la configuration électronique de l'ion d'un atome isolé, et cela en négligeant le potentiel de cœur. Or, en mécanique quantique toutes les fonctions d'onde décrivant les états électroniques doivent être orthogonales entre elles, donc toutes les fonctions d'ondes de valences doivent être orthogonales à celles du cœur, et c'est ce qui pose problème du côté numérique, car la fonction d'onde de valence va présenter une structure très compliquée (nodale). Pour cette raison, il serait donc plus efficace de remplacer le véritable ion par un potentiel effectif auquel est associé une fonction d'onde adoucie ou régulière (qui est sans nœuds), et cela est justifié : car le potentiel répulsif généré par les électrons du cœur est compensé par le potentiel attractif du noyau. Donc, il en résulte un potentiel ionique relativement faible dit «potentiel effectif» qui préserve les énergies propres du système, et les fonctions d'ondes associées sont appelées pseudo-fonctions.

Le concept de pseudo-potentiel a été introduit la première fois dans les années 1930 par Fermi, en suite, Hellmann utilise cette notion pour le calcul des niveaux énergétiques de métaux alcalins, ces premiers pseudo-potentiels sont qualifiés d'empiriques ; ce qui signifie qu'ils ne sont pas obtenus par calcul mathématiques mais ajustés pour reproduire au mieux des résultats expérimentaux de référence, ce qui suggère qu'il y a plusieurs façons de générer un pseudo-potentiel : soit paramétré pour reproduire des résultats expérimentaux ; et c'est le pseudo-potentiel empirique (comme dit précédemment), soit on se base sur une approche mathématique pour modifier la fonction d'onde électronique.

Cette approximation a de gros avantages en physique, non seulement elle réduit le nombre d'électrons du problème, en ne traitant explicitement que les électrons de valence, mais aussi elle réduit le nombre d'état de base, c'est à dire le nombre d'ondes planes utilisées pour décomposer la fonction d'onde, en sachant que ce nombre augmente avec l'augmentation des irrégularités de la fonction d'onde. Et ces avantages permettent d'économiser beaucoup sur le plan des ressources informatiques.

0.1.3 Théorème de Bloch

En partant du fait que le potentiel $V(\vec{r})$ est périodique, on peut dire qu'il y a une symétrie par translation.

Première formulation :

Les états propres de l'hamiltonien H peuvent être choisis de telle sorte qu'à chaque ψ on peut associer un vecteur d'onde $\vec{k}$ tel que :

$$\psi(\vec{r} + \vec{R}) = e^{-i\vec{k}\vec{R}}\psi(\vec{r})$$

Démonstration :

On cherche la solution l'équation de Schrödinger :

$$\left(-\frac{\hbar^2}{2m} \nabla^2 + V(\vec{r}) \right) \psi(\vec{r}) = \epsilon \psi(\vec{r}) \quad (1)$$

où $V(\vec{r})$ est un potentiel périodique :

$$V(\vec{r} + \vec{R}) = V(\vec{r})$$

et $\vec{R} = \sum_i^3 n_i \vec{a}_i$ avec $\vec{a}_i$ trois vecteurs linéairement indépendants.

Pour tenir compte de l'invariance par translation d'un vecteur $\vec{R}$ de $V(\vec{r})$ introduisons l'opérateur de translation $T_{\vec{R}}$ défini par :

$$T_{\vec{R}}\psi(\vec{r}) = \psi(\vec{r} + \vec{R}) \quad (2)$$

$\forall \vec{R}, \forall \psi \in$ l'espace de Hilbert.

Maintenant, si on applique deux translation, le resultat ne depend pas de l'ordre dans lequel sont appliquées les translations :

$$T_{\vec{R}}T'_{\vec{R}'} = \psi(\vec{r} + \vec{R} + \vec{R}')$$

on a donc :

$$T_{\vec{R}}T'_{\vec{R}'} = T'_{\vec{R}'}T_{\vec{R}} = T_{\vec{R}+\vec{R}'} \quad \forall(\vec{R}, \vec{R}') \quad (3)$$

L'opérateur de translation $T_{\vec{R}}$ commute d'autre part avec le hamiltonien H de (1) :

$$\begin{aligned} T_{\vec{R}} H \psi(\vec{r}) &= \left(-\frac{\hbar}{2m} \nabla^2 + V(\vec{r} + \vec{R}) \right) \psi(\vec{r} + \vec{R}) \\ &= \left(-\frac{\hbar}{2m} \nabla^2 + V(\vec{r}) \right) T_{\vec{R}} \psi(\vec{r}) \\ &= H T_{\vec{R}} \psi(\vec{r}) \end{aligned}$$

Donc :

$$[H, T_{\vec{R}}] = 0 \quad \forall \vec{R} \quad (4)$$

Le sens physique de cette relation est clair : Elle signifie que pour un état propre $\psi(\vec{r})$ de H donné, sa transformée $T_{\vec{R}} \psi(\vec{r})$ est aussi un état propre avec la même énergie. En effet l'équation (4) montre que pour toute fonction quelconque ϕ , on a $H T_{\vec{R}} \phi = T_{\vec{R}} H \phi$; choisissons pour ϕ un état propre de H , ψ satisfait : $H\psi = \epsilon\psi$, il vient alors :

$$H(T_{\vec{R}} \psi) = H\phi = T_{\vec{R}} H \psi = T_{\vec{R}}(\epsilon \psi) = \epsilon(T_{\vec{R}} \psi) \quad (5)$$

cela montre que $(T_{\vec{R}} \psi)$ est fonction propre de H avec la valeur propre ϵ , tout comme l'est ψ ! De toute évidence l'ensemble des $T_{\vec{R}}$ peut être muni d'une structure de groupe, au moyen de la relation :

$$T_{\vec{R}} T_{\vec{R}'} = T_{(\vec{R}+\vec{R}')}$$

L'inverse est défini comme :

$$T_{\vec{R}}^{-1} = T_{-\vec{R}}$$

ce groupe est ici de puissance N (translations discrètes) ; il est commutatif (abélien), puisque la somme $\vec{R} + \vec{R}'$ est commutative :

$$T_{\vec{R}} T_{\vec{R}'} = T_{(\vec{R}+\vec{R}')} = T_{(\vec{R}'+\vec{R})} = T_{(\vec{R}')} T_{\vec{R}} \implies [T_{\vec{R}}, T_{\vec{R}'}] = 0$$

Il en résulte que tous les $T_{\vec{R}}$ ont des vecteurs propres communs. Par ailleurs, les $T_{\vec{R}}$ sont des opérateurs unitaires puisqu'ils forment un sous-ensemble discret des translations continues ; parmi ces dernières existent les translations infinitésimales, qui ne sauraient être anti-unitaires. En conséquent :

$$T_{\vec{R}}^\dagger = T_{\vec{R}}^{-1} = T_{-\vec{R}}$$

Comme $T_{\vec{R}}$ et H commutent, alors ils ont des vecteurs propres en commun. Soit une

fonction propre $\psi(\vec{r})$ telle que les deux équations suivantes sont satisfaites :

$$H\psi(\vec{r}) = \epsilon\psi(\vec{r}) \quad (6)$$

$$T_{\vec{R}}\psi(\vec{r}) = \tau(\vec{R})\psi(\vec{r}) \forall \vec{R} \in \mathfrak{H} \quad (7)$$

En conséquence, de (7) et (2) on a :

$$T_{\vec{R}} T_{\vec{R}'} \psi(\vec{r}) = \tau(\vec{R})\tau(\vec{R}') \psi(\vec{r}) \quad (8)$$

$$T_{(\vec{R}+\vec{R}')} \psi(\vec{r}) = \tau(\vec{R} + \vec{R}') \psi(\vec{r}) \quad (9)$$

Les equations (8) et (9) sont vraies $\forall \psi$ également fonction propre de l'hamiltonien H .
Donc il en résulte :

$$\tau(\vec{R})\tau(\vec{R}') = \tau(\vec{R} + \vec{R}') \quad (10)$$

τ est donc une fonction exponentielle -d'ailleurs, en vertu de l'unitarité des $T_{\vec{R}}$, toutes leurs valeurs propres sont de la forme $e^{(i \times \text{phase réelle})}$ - ce que l'on choisit d'écrire plus précisément, pour la commodité :

$$\tau(\vec{R}) = e^{-i\phi(\vec{R})} \quad (11)$$

où ϕ est une fonction à valeurs réelles, (10) s'écrit alors :

$$\phi(\vec{R} + \vec{R}') = \phi(\vec{R}) + \phi(\vec{R}') \quad (12)$$

Cette relation dit que ϕ est une forme linéaire de $\vec{R} \forall \vec{R}$, une telle forme linéaire peut toujours s'écrire en produit scalaire, $\phi(\vec{R})$ est donc de la forme $\vec{k} \cdot \vec{R}$, et en définitive

$$\tau(\vec{R}) = e^{i\vec{k} \cdot \vec{R}} \quad (13)$$

Compte tenu de ce résultat, la relation (7) s'écrit maintenant , $\forall \psi$ propre de H et des $T_{\vec{R}}$:

$$\begin{aligned} T_{\vec{R}} \psi(\vec{r}) &= e^{i\vec{k} \cdot \vec{R}} \psi(\vec{r}) \\ \psi(\vec{r} + \vec{R}) &= e^{i\vec{k} \cdot \vec{R}} \psi(\vec{r}) \end{aligned} \quad (14)$$

et ceci n'est autre que la première formulation du théorème de Bloch.

Deuxième formulation :

Les états propres de H peuvent s'écrire comme la combinaison d'une onde plane et

d'une fonction arbitraire qui possède la périodicité du réseau

$$\psi_n(\vec{r}) = e^{i\vec{k}\vec{r}} u_n(\vec{r}) \quad (15)$$

Démonstration

Considérons la fonction d'onde d'un électron libre :

$$\psi_{\vec{k}}(\vec{r}) = \frac{1}{\sqrt{V}} e^{i\vec{k}\vec{r}} \quad (16)$$

On peut dire qu'elle satisfait le théorème de Bloch. On doit s'y attendre car on peut considérer qu'un électron libre se déplace dans un potentiel périodique qui est partout nul. Il est utile d'introduire une fonction d'onde dont la forme rappelle celle d'une onde plane en posant :

$$\psi_n(\vec{r}) = e^{i\vec{k}\vec{r}} u_n(\vec{r}) \quad (17)$$

La relation (17) est compatible avec le théorème de Bloch si $u_n(\vec{r})$ possède la périodicité du réseau, soit :

$$u_n(\vec{r}) = u_n(\vec{r} + \vec{R}) \quad (18)$$

On peut le démontrer comme suit :

$$\begin{aligned} \psi(\vec{r} + \vec{R}) &= e^{(i\vec{k}\vec{r} + i\vec{k}\vec{R})} u_n(\vec{R} + \vec{r}) \\ &= e^{i\vec{k}\vec{r}} e^{i\vec{k}\vec{R}} u_n(\vec{R} + \vec{r}) \\ &= e^{i\vec{k}\vec{R}} \psi(\vec{r}) \\ &= e^{i\vec{k}\vec{R}} e^{i\vec{k}\vec{r}} u_n(\vec{r}) \end{aligned}$$

En simplifiant :

$$u_n(\vec{R} + \vec{r}) = u_n(\vec{r})$$

Et de cette façon on a démontré la seconde formulation du théorème de Bloch.

0.1.4 Modèle de Krönig-Penney

Le modèle de Krönig-Penney est utilisé pour remplacer le potentiel périodique d'un réseau cristallin à une dimension, par un puits carré à chaque site du réseau. Dans ce modèle, on assume que $V(x)$ est nul partout excepté sur les sites atomiques (figure1) :

$$V(x) = \begin{cases} -V_0 & \text{si } na < x \leq na + b \\ 0 & na + b < x \leq (n+1)a \end{cases} \quad V_0 > 0, \quad n \in \mathbb{Z}$$

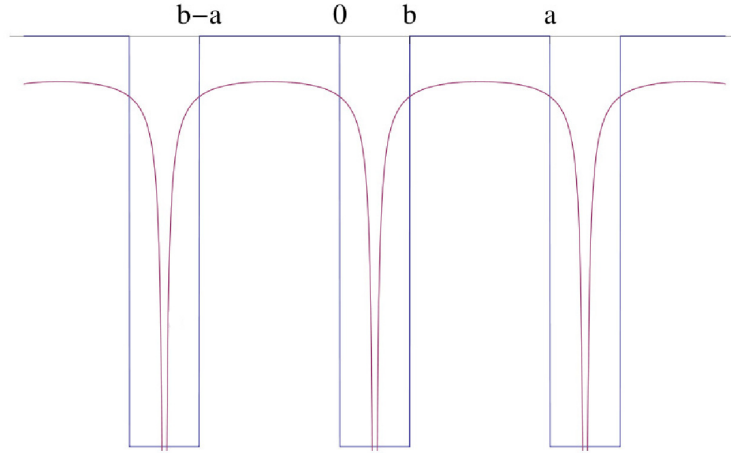


FIGURE 1 – Le potentiel périodique ressenti par un électron dans un réseau unidimensionnel d’atomes (ligne mauve) et la schématisation avec des puits de potentiel carrés (ligne bleue).

Nous cherchons les états liés ($E < 0$) du système :

Une première façon de procéder consisterait à écrire la solution générique de l’équation de Schrödinger dans chaque région, puis, les raccorder avec les conditions aux limites adéquates.

Pour un réseau infini, il est clair que cette procédure n’est pas réalisable. Cependant, l’implémenter numériquement est faisable en laissant le nombre de puits comme un paramètre modifiable à la guise de l’expérimentateur, et c’est de cette façon que nous allons procéder lors de notre étude numérique. Mais en ce qui concerne le modèle de Krönig-Penney la meilleure solution est d’utiliser le théorème de Bloch.

Solution de l’équation de Schrödinger

Considérons la maille entre $x = b - a$ et $x = b$: Nous pouvons identifier deux régions à potentiel constant (I) : $b - a < x < 0$ et (II) : $0 < x < b$. La solution générique dans la région (I) s’écrit :

$$\psi_I = A \cosh(\lambda x) + B \sinh(\lambda x) \quad (19)$$

$$\lambda = \sqrt{\frac{-2mE}{\hbar^2}} \quad (20)$$

Dans la région (II) :

$$\psi_{II} = C \cos(\lambda'x) + D \sin(\lambda'x) \quad (21)$$

$$\lambda' = \sqrt{\frac{-2m(V_0 + E)}{\hbar^2}} \quad (22)$$

Les deux conditions aux bords en $x = 0$:

$$\psi_I(0) = \psi_{II}(0)$$

$$\psi'_I(0) = \psi'_{II}(0)$$

impliquant immédiatement que :

$$A = C$$

$$B = \frac{\lambda'}{\lambda} D$$

Les fonctions d'onde sont donc de la forme

$$\psi_I = C \cosh(\lambda x) + \frac{\lambda'}{\lambda} D \sinh(\lambda x)$$

$$\psi_{II} = C \cos(\lambda x) + D \sin(\lambda' x)$$

Il faut également raccorder la fonction d'onde dans la maille $[b-a, b]$ aux fonctions d'onde dans la maille $[b-2a, b-a]$ et $[b, b+a]$; pour cela, on va utiliser le théorème de Bloch (Eq (15)). Pour fixer les idées, considérons la maille $[b, b+a]$ et appelons *III* et *I* les intervalles $[b, a]$ et $[a, b+a]$ respectivement. On devrait écrire les conditions de raccord entre les mailles $[b-a, b]$ et $[b, b+a]$ de la façon suivante :

$$\psi_{II}(b) = \psi_{III}(b)$$

$$\psi'_{II}(b) = \psi'_{III}(b)$$

Mais le théorème de Bloch nous dit que la fonction d'onde doit être quasi-périodique

$$\psi_k(b) = e^{ika} \psi_k(b-a)$$

$$\psi'_k(b) = e^{ika} \psi'_k(b-a)$$

Par conséquence :

$$\psi_{III}(b) = e^{ika} \psi_I(b-a)$$

$$\psi'_{III}(b) = e^{ika} \psi'_I(b-a)$$

et les conditions de raccord avec la maille $[b, b + a]$ peuvent être écrites comme :

$$\psi_{II}(b) = e^{ika}\psi_I(b - a) \quad (23)$$

$$\psi'_{II}(b) = e^{ika}\psi'_I(b - a) \quad (24)$$

En termes des coefficients C et D, ces deux équations représentent un système d'équations linéaires à deux inconnus :

$$\begin{bmatrix} \cos(\lambda'b) - e^{ika}\cosh\lambda(b-a) & \sin(\lambda'b) - \frac{\lambda'}{\lambda}e^{ika}\sinh\lambda(b-a) \\ -\lambda'\sin(\lambda'b) - \lambda e^{ika}\sinh\lambda(b-a) & \lambda'\cos(\lambda'b) - e^{ika}\cosh\lambda(b-a) \end{bmatrix} \begin{bmatrix} C \\ D \end{bmatrix} = 0$$

On aura des solution seulement si :

$$\begin{vmatrix} \cos(\lambda'b) - e^{ika}\cosh\lambda(b-a) & \sin(\lambda'b) - \frac{\lambda'}{\lambda}e^{ika}\sinh\lambda(b-a) \\ -\lambda'\sin(\lambda'b) - \lambda e^{ika}\sinh\lambda(b-a) & \lambda'\cos(\lambda'b) - e^{ika}\cosh\lambda(b-a) \end{vmatrix} = 0$$

En utilisant l'identité $\cosh^2x - \sinh^2x = 1$, on trouve la condition de quantification suivante :

$$\cos(\lambda'b)\cosh\lambda(b-a) - \frac{\lambda^2 - \lambda'^2}{2\lambda\lambda'}\sin(\lambda'b)\sinh\lambda(b-a) = \cos ka \quad (25)$$

On voit qu'on obtient une équation de quantification différente pour chaque valeur de k .

Pour simplifier l'équation (25), considérons la limite du potentiel delta (peigne de Dirac). On impose :

$$V_0 = \alpha = cst$$

et on prend la limite pour $b \rightarrow 0$:

$$\begin{aligned} \lambda'b &= \sqrt{\frac{2m(V_0 + E)}{\hbar^2}} b \\ &= \sqrt{\frac{2m(\alpha/b + E)}{\hbar^2}} b \\ &\simeq \sqrt{\frac{2m\alpha b}{\hbar^2}} \quad (b \rightarrow 0) \end{aligned}$$

et donc

$$\begin{aligned} \sin\lambda'b &\simeq \sqrt{\frac{2m\alpha b}{\hbar^2}} \\ \cos\lambda'b &\simeq 1 \end{aligned}$$

En outre, on a

$$\lambda(b - a) \simeq -\lambda a \quad (b \rightarrow 0)$$

et

$$\begin{aligned} \frac{\lambda^2 - \lambda'^2}{\lambda\lambda'} &= \frac{-E - (V_0 + E)}{\sqrt{-E(V_0 + E)}} \\ &= \frac{-V_0 - 2E}{\sqrt{-E(V_0 + E)}} = \frac{-\alpha/b - 2E}{\sqrt{-E(\alpha/b + E)}} \\ &\simeq \frac{-\alpha/b}{\sqrt{-E\alpha/b}} \quad (b \rightarrow 0) \\ &= -\sqrt{\frac{\alpha}{-bE}} \end{aligned}$$

La condition de quantification (25) devient alors

$$\cosh\lambda a - \frac{m\alpha}{\hbar^2\lambda} \sinh\lambda a = \cos ka \quad (26)$$

Cette équation nécessite une résolution numérique, donc on va se contenter de décrire son comportement d'une manière qualitative.

Discussion

Si nous observons l'équation (26), nous remarquons que le deuxième membre de cette équation est compris dans l'intervalle $[-1, 1]$. Les valeurs de λ (et donc de l'énergie : $\lambda = \sqrt{\frac{-2mE}{\hbar^2}}$) pour lesquelles le premier membre est en dehors de cette intervalle sont donc interdites (ces valeurs de l'énergies forment le gap). Le vecteur d'onde k prendra donc les valeurs entre $[\frac{-n\pi}{a}, \frac{n\pi}{a}]$ avec n un entier positif, pour chaque valeur de k nous aurons une valeur de l'énergie E_k , et toutes les énergies appartenant à un intervalle $[E_{\frac{-n\pi}{a}}, E_{\frac{n\pi}{a}}]$ (pour un n particulier) sont permises et forment donc une bande.

0.1.5 Conducteurs, semiconducteurs et isolants

Les propriétés électriques d'un matériau sont fonctions du remplissage électronique des différentes bandes permises. La conduction électrique résulte du déplacement des électrons à l'intérieur de chaque bande, sous l'application d'un champ électrique. Considérons à présent une bande d'énergie vide, il est évident de part le fait qu'elle est vide, qu'elle ne participe pas à la formation d'un courant électrique. Il en est de même pour une bande complètement pleine. En effet, un électron ne peut se déplacer que si il existe une place

vide (un trou) dans la bande d'énergie où il se trouve. Ainsi, un matériau dont les bandes d'énergies sont vides ou complètement pleines est un isolant. Un tel comportement est obtenu pour les énergies de gap supérieur à 9eV , car pour de telles énergies l'agitation thermique à 300K , ne peut pas faire passer les électrons de la bande de valence à la bande de conduction.

Un semiconducteur est un isolant pour une température de 0K . Cependant, ce type de matériau ayant une énergie de gap inférieur à celle d'un isolant $\sim 1\text{eV}$, aura de par l'agitation thermique ($T = 300\text{K}$), une bande de conduction légèrement peuplée, et une bande de valence légèrement dépeuplée par le fait du déplacement des électrons à la bande de conduction. On en déduit donc que la conductivité d'un semiconducteur est faible et elle est dépendante de la température : Quand la température augmente la conductivité du semiconducteur augmente aussi.

Pour un conducteur, la bande de conduction est partiellement remplie, et celle de valence est partiellement vide. Dans cette situation, il est très facile pour les électrons de valence de passer aux niveaux d'énergies supérieurs, et ainsi participer à la conduction électrique sous l'effet d'un champ électrique. Contrairement aux semiconducteurs, la conductivité d'un métal diminue avec l'augmentation de la température, et cela à cause de l'agitation thermique qui gêne le déplacement des électrons, en d'autres termes l'augmentation de la température augmente la résistivité du conducteur.

Comme cité auparavant, la conductivité d'un semiconducteur est faible, pour l'améliorer on utilise la procédure de dopage.

Dopage des semiconducteurs

Nous allons ici donner un bref aperçu sur l'opération de dopage des semiconducteurs³ pour l'intérêt qu'elle suscite que ça soit en physique ou en ingénierie.

L'opération de dopage d'un milieu semiconducteur simple ou composé est basée sur une introduction (implantation) d'atomes impuretés de nature chimique trivalente⁴ (B, Al, Ga, ...) ou pentavalente⁵ (N, As, Sb, P, ...). Après dopage les propriétés physiques du semiconducteur sont plus au moins modifiées par la présence de ces atomes dopants. Par définition, le dopage d'un semiconducteur avec des atomes trivalents va donner lieu à un semiconducteur de type P (positif), le semiconducteur résultant d'un dopage avec les atomes pentavalents est dit de type N (négatif).

L'objectif principal recherché à travers ces opérations de dopage de ces milieux semi-

3. On appelle aussi le semiconducteur dopé un semiconducteur extrinsèque. Inversement un semiconducteur qui n'est pas dopé s'appelle intrinsèque

4. Il a trois électrons de valence

5. Il a cinq électrons de valence

conducteurs est l'amélioration de leurs propriétés électriques à travers l'accroissement de leur conductivité électrique (sans pour autant atteindre celle des milieux métalliques c'est à dire : tout en préservant le caractère semiconducteur du matériau). Par définition, les atomes trivalents sont appelés atomes « accepteurs », l'ionisation de chacun de ces atomes à l'intérieur du semiconducteur est réalisé à travers la capture d'un électron de la bande de valence, à l'opposé les atomes pentavalents dit « donneurs » vont manifester une ionisation à l'intérieur du semiconducteur à travers une perte d'un de leur électrons de valence (par atome dopant).

Le recourt à des dopages avec de faible teneurs d'atomes impureté est appliqué de manière à conserver la structure cristalline du matériau semiconducteur de départ (La matrice cristalline hôte), et à s'assurer d'une disposition atomique substitutionnelle⁶ de ces atomes impureté à l'intérieur du semiconducteur hôte.

Les propriétés physico-chimiques des semiconducteurs dopés sont très dépendantes de la nature chimique des atomes dopants entrants dans leur composition. En général, le dopage des milieux semiconducteurs avec des atomes donneurs va générer des électrons libres supplémentaire dans la bande de conduction. A l'opposé, un dopage avec des atomes accepteurs est à l'origine d'un supplément de trous mobiles dans la bande de valence.

A température ambiante $T = 300K$, les atomes de dopages sont supposées entièrement ionisés. La neutralité électrique globale de l'échantillon semiconducteur est ainsi assurée par une égalité entre les densités de charges négatives (les électrons) et les densités de charges positives (les trous mobiles) des différents constituants : les électrons libres des semiconducteurs + les atomes dopants ionisés + les trous mobiles de la bande de valence.

A $T = 300K$ il y a une ionisation complète des atomes de dopage et les densités de charges intrinsèques sont très petites devant les densités de charges provenant des atomes dopants : $\tilde{n}_i \ll \tilde{n}_d$. Cette ionisation complète des atomes de dopage à température ambiante va donner lieu aux densités de porteurs de charges suivantes :

Cas du dopage N

on appelle $\tilde{n}$, $\tilde{n}_i$, $\tilde{n}_d$ les densités de charge électroniques (n : pour négatives) après dopage, intrinsèque et dopant donneur du semiconducteur respectivement, et $\tilde{P}$, $\tilde{P}_i$ sont les densités de charges des trous mobiles totales et intrinsèques respectivement. On a :

$$\begin{cases} \tilde{n} &= \tilde{n}_i + \tilde{n}_d \dots \dots \dots (\text{Bande de conduction}) \\ \tilde{P} &= \tilde{P}_i (= \tilde{n}_i) \dots \dots \dots (\text{Bande de valence}) \end{cases}$$

6. le dopage par substitution est un dopage qui a lieu en remplaçant l'un des atome du semiconducteur par l'atome dopant, on l'utilise souvent quand la taille de l'atome dopant est grande

avec

$$\tilde{n}_i \ll \tilde{n}_d = \tilde{N}_d = \frac{\text{nombre de charge donneur}}{\text{volume du semiconducteur}}$$

On remarque que les densités de charges électroniques et des trous mobiles dans un semiconducteur intrinsèque c'est à dire avant dopage sont égales d'où la neutralité du semiconducteur.

Cas du dopage P

$$\begin{cases} \tilde{n} &= \tilde{n}_i \dots\dots\dots (\text{Bande de conduction}) \\ \tilde{P} &= \tilde{P}_i + \tilde{P}_p (= \tilde{n}_i) \dots\dots\dots (\text{Bande de valence}) \end{cases}$$

tel que : $\tilde{P}$, $\tilde{P}_i$ et $\tilde{P}_a$ sont respectivement : les densités de charge du semiconducteur dopé P, les densités de charges positives du semiconducteur intrinsèque et les densité de charges des atomes dopants trivalents : accepteurs.

avec

$$\tilde{P}_i \ll \tilde{P}_a = \tilde{N}_a = \frac{\text{nombre de charge accepteur}}{\text{volume du semiconducteur}}$$

En pratique, les densités intrinsèques $\tilde{n}_i$ et $\tilde{P}_i$ du semiconducteur avant dopage sont largement inférieurs aux densités des électrons libres et trous mobiles générés par les atomes de dopage donneurs et accepteurs respectivement.

$$(\tilde{n}_i \text{ et } \tilde{P}_i) \ll (\tilde{n}_d \text{ et } \tilde{P}_a)$$

Exemple :

$$\begin{aligned} \tilde{n}_i^{SI} &= 1.5 \times 10^{16} [elec/m^3]_{SI} \\ \tilde{N}_d^{(P)} &= 10^{19} \text{ à } 10^{23} [atm/m^3]_{SI} \end{aligned}$$

SI : Pour système international.

Dans le cas d'un dopage avec les atomes donneurs, la densité des électrons libres dans la bande de conduction est largement supérieur à celle des trous mobiles de la bande de valence. Dans ces semiconducteurs de type «N» les électrons libres sont ainsi considérés comme des porteurs de charges «Majoritaires» et les trous mobiles comme des porteurs de charges «Minoritaires» .

$$\text{semiconducteur «N»} \begin{cases} \tilde{n} &= \tilde{n}_i + \tilde{n}_d \approx \tilde{n}_d & (\text{Majoritaire}) \\ \tilde{P} &= \tilde{P}_i \ll \tilde{n}_d & (\text{Minoritaire}) \end{cases}$$

avec $\tilde{n}_i \ll \tilde{n}_d$ à $T = 300K$.

Dans le cas contraire, en dopant avec des atomes accepteurs, ce sont les trous mobiles de la bande de valence qui sont considérés comme des porteurs de charge «Majoritaires» et les électrons libres de la bande de valence comme des porteurs de charge «Minoritaires».

$$\text{semiconducteur «P»} \begin{cases} \tilde{n} &= \tilde{n}_i \ll \tilde{P}_a & (\text{Minoritaire}) \\ \tilde{P} &= \tilde{P}_i + \tilde{P}_a \approx \tilde{P}_a & (\text{Majoritaire}) \end{cases}$$

avec $\tilde{P}_i \ll \tilde{P}_a$.

En générale, le comportement électrique des semiconducteurs de type «N» fortement dopés ($\tilde{N}_d \gg$) est très proche de celui des métaux en raison des densités de charges uniformes (constantes) les caractérisant $\tilde{n} = \tilde{n}_d = \text{constante}$, et pour le fait que les densités de charges de ces milieux sont indépendantes de la température.

Les semiconducteurs résultant d'un dopage avec des atomes donneurs sont dit de type «négatifs» ou «N» en raison de la charge électrique négative porté par leur porteurs de charge majoritaires (c'est à dire : porté par les électrons libre de la bande de conduction). Du point de vu de la neutralité globale d'un semiconducteur extrinsèque hybride, en d'autres termes : dopé à la fois avec des atomes donneurs « N_d » et « N_a » atomes accepteurs, le matériau semiconducteur reste neutre avant et après dopage, on peut décrire cette neutralité comme suit :

Avant dopage

$$\tilde{n}_i(-|e|) + \tilde{P}_i(+|e|) = 0 \implies \text{échantillon neutre avant le dopage.}$$

Après dopage hybride

$$\tilde{n}(-|e|) + \tilde{P}(+|e|) + \tilde{N}_a^-(-|e|) + \tilde{N}_+^d(+|e|) = 0 \implies \text{échantillon neutre après dopage.}$$

Le dopage d'un semiconducteur avec des densités égales en atomes donneurs et accepteurs $\tilde{N}_d = \tilde{N}_a$ va donner lieu à un comportement électrique similaire à celui du semiconducteur intrinsèque. Ce type particulier de matériau est appelé «semi-isolant», à basse température, il est décrit par l'équation de neutralité : $\tilde{n} = \tilde{P}$

0.2 Réseaux irréguliers : La localisation d'Anderson

En 1958, Philipe W.Anderson[3] a publié un article où il a discuté du comportement des électrons dans un cristal imparfait. Après son étude sur les semiconducteurs, il a prédit

l'existence d'un régime localisé correspondant à une absence totale de diffusion quand le cristal contient un grand nombre d'impuretés et que le désordre est suffisamment fort. Ce phénomène porte son nom et s'appelle «localisation d'Anderson» ou «localisation forte».

Pour mieux comprendre la localisation d'Anderson on va prendre le cas opposé, celui d'un cristal parfait, et comme exemple on va choisir le cas d'un métal idéal. Dans ce cas là, la conduction électrique est assurée par un gaz d'électron quasi-libres, mis en mouvement par un champ électrique extérieur. Mais dans le cas réel, la plupart des cristaux portent des défauts et contiennent des impuretés qui peuvent modifier radicalement leur comportement électrique. En fait, les électrons vont subir de nombreuses collisions sur les imperfections du cristal, acquérant un mouvement global diffusif.

Dans ses travaux traitant des électrons, Drude montre dans son modèle⁷ que la conductivité d'un métal est déterminée par le libre parcours moyen des électrons⁸. Ainsi, si le libre parcours moyen est grand, cela veut dire que l'électron peut se propager rapidement sur des grandes distances, par conséquent, on aura une conductivité élevée. Maintenant, en augmentant le degré du désordre dans l'échantillon, le libre parcours moyen diminue, et la conductivité diminue à son tour. Parmi les failles de ce modèle, le fait de considérer les électrons comme des particules classiques et non pas comme des particules quantiques décrites par une fonction d'onde. Or, quand on parle de fonction d'onde, on parle aussi d'interférences. Dans notre cas, la partie de la fonction d'onde électronique diffusée par la présence de l'impureté peut interférer avec la partie non diffusée, et cela peut affecter considérablement le transport.

En mécanique quantique, pour décrire des ondes de matière tel que les électrons, on utilise une longueur caractéristique qui est la longueur d'onde de de Broglie⁹, plus particulièrement la longueur d'onde thermique de de Broglie, étant donné qu'on décrit un gaz d'électron¹⁰. C'est le rapport entre le libre parcours moyen et la longueur d'onde de de Broglie qui va déterminer la nature du transport électrique en présence de désordre. Si le libre parcours moyen est supérieur à la longueur d'onde de de Broglie, la phase accumulée

7. Le modèle classique de Drude traite la conductivité électrique dans les métaux. Dans ce modèle, il considère les électrons dans un métal comme un gaz de particules libres et sans interaction mutuelle, remplissant uniformément tous les espaces inter-atomiques, et surtout, il considère les électrons comme des particules classiques régies par la statistique de Maxwell-Boltzman.

8. Le libre parcours moyen est une longueur caractéristique qui représente la distance moyenne parcourue par l'électron entre deux chocs successifs

9. Louis de Broglie énonce que toute particule physique dotée d'une quantité de mouvement, a une onde associée de longueur d'onde appelée longueur d'onde de de Broglie.

10. La longueur d'onde thermique de de Broglie est une grandeur statistique qui représente la longueur d'onde de de Broglie moyenne d'un gaz porté à une certaine température. Elle caractérise l'étalement spatiale de la particule associée et fait le lien entre la mécanique classique et la mécanique quantique. En effet le caractère quantique commence à être important lorsque la longueur d'onde de de Broglie est comparable aux autres longueurs caractéristiques du système tel que le libre parcours moyen ou le volume du système

entre deux diffusions successives est très grande devant 2π , et les effets d'interférences se moyennent à zéro, on retrouve donc le comportement classique du modèle de Drude : C'est le cas par exemple d'un métal usuel où le libre parcours moyen est de l'ordre de 100 nm, tandis que la longueur d'onde de de Broglie est d'environ 1 nm. Quand à Anderson, lui, il s'est demandé si le modèle de Drude restait valable si le désordre était élevé. Il a alors construit le modèle du cristal désordonné.

Dans son modèle, Anderson a utilisé une approche par liaisons fortes (plus adéquate pour décrire des états localisés) à un électron pour étudier le cristal :

On commence par le hamiltonien suivant :

$$H = \sum_i \epsilon_i |i\rangle\langle i| + \sum_{i \neq j} t_{ij} |i\rangle\langle j|$$

ou $|i\rangle$ est le vecteur d'état sur chaque site i , ϵ_i est le niveau d'énergie de l'électron sur chaque site i , et t_{ij} est l'intégrale de transfert¹¹ entre deux sites i et j . Dans le modèle d'Anderson, l'intégrale de transfert est non nulle uniquement pour les premiers proches voisins, et elles sont supposées être égales, c-à-d : $t_{ij} = \tilde{t}$, et cela est dû au fait que l'on a pris la même distance entre chaque site, à l'opposé les ϵ_i sont choisis pour qu'ils soient aléatoires sur chaque site $|i\rangle$, et variant dans l'intervalle $-W/2 \leq \epsilon_i \leq W/2$, où W est la valeur du potentiel.

Si on développe la fonction d'onde en terme du modèle de liaisons fortes :

$$|\psi\rangle = \sum_i a_i |i\rangle$$

où $|i\rangle$ est l'orbitale atomique centrée en i .

Admettons que ψ est une fonction propre qui satisfait $H\psi = E\psi$, on peut alors obtenir

11. Cette intégrale de transfert représente le couplage par effet tunnel entre les sites

l'équation pour l'amplitude a_i :

$$\begin{aligned}
H\psi &= E\psi \\
\left(\sum_i \epsilon_i |i\rangle\langle i| + \sum_{i \neq j} t_{ij} |i\rangle\langle j| \right) \psi &= E\psi \\
\left(\epsilon_i |i\rangle\langle i| + \sum_j t_{ij} |i\rangle\langle j| \right) a_i |i\rangle &= E a_i |i\rangle \\
\epsilon_i a_i |i\rangle + \sum_j t_{ij} a_j |i\rangle &= E a_i |i\rangle \\
\epsilon_i a_i |i\rangle + \sum_j t_{ij} a_j \delta_{ij} |i\rangle &= E a_i |i\rangle \\
\epsilon_i a_i |i\rangle + \sum_j t_{ij} a_j |i\rangle &= E a_i |i\rangle
\end{aligned}$$

en multipliant par $\langle i|$:

$$\epsilon_i a_i \langle i|i\rangle + \sum_j t_{ij} a_j \langle i|i\rangle = E a_i \langle i|i\rangle$$

$$E a_i = \epsilon_i a_i + \sum_j t_{ij} a_j \quad (27)$$

et cela donne la solution d'un état stationnaire.

Maintenant, si on considère l'équation du mouvement dépendante du temps, on aura :

$$-i\hbar \frac{d}{dt} a_i = \epsilon_i a_i + \sum_j t_{ij} a_j \quad (28)$$

En observant cette équation, on peut définir ce que c'est qu'un état localisé.

Supposons qu'à $t = 0$, un électron se place sur le site i , et que $a_i(t = 0) = 1$, mais $a_j = 0$ pour $i \neq j$. Grâce à ces conditions initiales, on va déterminer l'évolution du système décrit par l'équation(28).

Examinons $a_i(t)$ dans la limite des temps infinis. Si $a_i(t \rightarrow \infty) = 0$, en d'autres termes l'électron ne se trouve pas sur le site de départ (on a posé comme condition initiale $a_i(t = 0) = 1$), alors l'électron est dans un état étendu. Maintenant, si $a_i(t \rightarrow \infty)$ est égale à une valeur finie, cela veut dire que l'électron est confiné dans la région des plus proches voisins. Il est donc dans un état localisé. Il conclut alors que pour un désordre suffisamment élevé, le mouvement des électrons est complètement stoppé, et la

conductivité du matériau s'annule et il devient isolant.

Deuxième partie

Approche numérique

Dans cette section nous avons essayé de mettre en évidence les propriétés étudiées précédemment pour une particule dans un solide. Pour rendre l'étude évidente, nous avons choisi que la particule en question va se mouvoir dans un potentiel périodique. De telles structures se rencontrent par exemple lors de l'étude d'une molécule linéaire, formée de n atomes (ou groupe d'atomes) identiques et régulièrement espacés. On les rencontre également en physique du solide, lorsqu'on prend un modèle à une dimension pour comprendre l'allure des niveaux d'énergie d'un électron dans un cristal, ce qui est le cadre de l'étude actuelle. Alors pour simplifier les choses, on choisira un potentiel formé de puits carrés. Ainsi nous allons voir ce que la structure de bande est réellement obtenue, nous allons voir aussi quelle est la forme des fonctions d'ondes. Puis, nous allons observer les changements survenus par l'introduction d'une impureté dans le réseau régulier de départ. Ensuite, nous allons essayer de simuler un cristal amorphe en rendant le réseau irrégulier par le changement de la distance qui sépare les puits de potentiel, et cela en utilisant des nombres aléatoires.

Pour la taille du réseau, nous savons que le potentiel $V(x)$ est réellement considéré comme périodique que quand $n \rightarrow \infty$ et l'on s'attend à ce que les propriétés de la particule soient pratiquement les mêmes que si $V(x)$ était réellement périodique. Cependant, d'un point de vue physique, la limite n infini n'est jamais réalisée, et dans notre étude numérique nous allons prendre $n = 10$.

Pour la réalisation de cette étude, nous avons utilisé deux méthodes pour s'assurer de la justesse des résultats.

0.3 Résolution de l'équation de Schrödinger à une dimension pour un électron dans un réseau périodique de puits carrés de potentiel

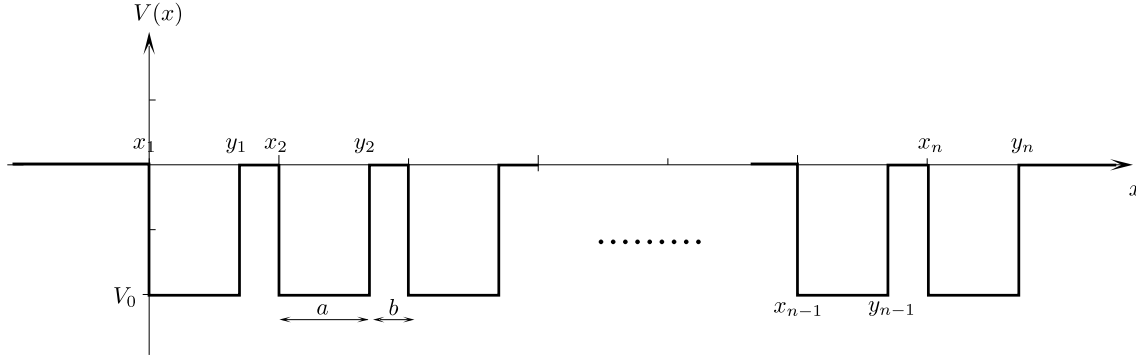
Nous résolvons numériquement l'équation de Schrödinger

$$\left(\frac{-\hbar^2}{2m} \frac{d^2}{dx^2} + V(x) \right) \psi(x) = E\psi(x) \quad (29)$$

pour une particule de masse m dans un potentiel à n puits défini par :

$$V(x) = \begin{cases} V_0 & \text{si } x \in [x_i, y_i[\\ 0 & \text{ailleurs} \end{cases} \quad \text{avec} \quad \begin{cases} x_i = (a+b)(i-1) \\ y_i = x_i + a \\ (i = 1, \dots, n) \end{cases} \quad (30)$$

L'allure du potentiel est :



Comme dit précédemment, nous présentons ici 2 méthodes plus ou moins différentes pour trouver les valeurs propres ainsi que les fonctions propres pour les états liés (c.à.d pour $V_0 < E < 0$, ou encore $\psi(x) \rightarrow 0$ lorsque $x \rightarrow \pm\infty$), et nous allons comparer les résultats obtenus avec ces méthodes. Avant d'exposer les méthodes ré-écrivons les équations précédentes pour y faire apparaître uniquement les variables sans dimensions qui sont plus adaptées à l'implémentation sur machine¹². On pose :

$$t = x/(a+b) \quad ; \quad x'_i = x_i/(a+b) \quad ; \quad y'_i = y_i/(a+b)$$

$$\begin{aligned} \gamma &= \sqrt{-2mV_0(a+b)^2/\hbar^2} \\ \xi &= \sqrt{E/V_0} \quad (> 0 \text{ pour les états liés ; si } E : V_0 \rightarrow 0 \text{ alors } \xi : 1 \rightarrow 0) \\ U(t) &= V(x)/V_0 \quad \left(U(t) \text{ compris entre } 1 \text{ et } 0 \right) \\ \phi(t) &= \sqrt{a+b} \psi(x) \quad (\text{pour avoir } \phi(t)^2 dt = \psi(x)^2 dx) \end{aligned} \quad (31)$$

Les équations (29) et (30) s'écrivent alors :

$$\phi''(t) = \gamma^2 \left(\xi^2 - U(t) \right) \phi(t) \quad (32)$$

12. Dans ce travail les programmes sont faits en C

$$U(t) = \begin{cases} 1 & \text{si } t \in [x'_i, y'_i[\quad (i = 1, \dots, n) \\ 0 & \text{ailleurs} \end{cases} \quad \text{avec} \quad \begin{cases} x'_i = x_i/(a+b) \\ y'_i = y_i/(a+b) \\ (i = 1, \dots, n) \end{cases}$$

Pour cette méthode, on résoud l'équation (32) dans chaque région, puis on impose que la fonction $\phi(t)$ et sa dérivé soient continues partout (en particulier aux points x'_i et y'_i , $i = 1, \dots, n$), et on impose aussi que $\phi(t) \rightarrow 0$ quand $t \rightarrow \pm\infty$ pour obtenir les états liés. Uniquement certaines valeurs de E (c'est à dire de ξ puisque $E = V_0\xi^2$) donnent des solutions acceptables physiquement : ce sont les valeurs propres, et on ne peut les déterminer que numériquement. Dans les régions $t \in [x'_i, y'_i[$ ($i = 1, \dots, n$), l'équation (32) s'écrit

$$\phi''(t) = \gamma^2(\xi^2 - 1)\phi(t) \quad \longrightarrow \quad \begin{cases} \text{solution générale :} \\ \phi(t) = A \cos(\beta t) + B \sin(\beta t) \\ \text{où } \beta \equiv +\gamma\sqrt{1 - \xi^2} \end{cases}$$

et dans les autres régions, on a :

$$\phi''(t) = \gamma^2 \xi^2 \phi(t) \quad \longrightarrow \quad \begin{cases} \text{solution générale :} \\ \phi(t) = C e^{\alpha t} + D e^{-\alpha t} \\ \text{où } \alpha \equiv +\gamma\xi \end{cases}$$

De manière plus précise, voici les résultats dans chaque région (on a posé ici $x'_1 = 0$) :

région $t \in]-\infty, x'_1[:$ $\phi(t) = A_0 e^{\alpha t}$ $\phi'(t) = \alpha A_0 e^{\alpha t}$	région $t \in [y'_n, +\infty[:$ $\phi(t) = A_{n+1} e^{-\alpha t}$ $\phi'(t) = -\alpha A_{n+1} e^{-\alpha t}$
région $t \in [x'_i, y'_i[\quad (i = 1, \dots, n) :$ $\phi(t) = A_i \cos(\beta t) + B_i \sin(\beta t)$ $\phi'(t) = \beta [-A_i \sin(\beta t) + B_i \cos(\beta t)]$	région $t \in [y'_{i-1}, x'_i[\quad (i = 2, \dots, n) :$ $\phi(t) = C_{i-1} e^{\alpha t} + D_{i-1} e^{-\alpha t}$ $\phi'(t) = \alpha [C_{i-1} e^{\alpha t} - D_{i-1} e^{-\alpha t}]$

Les A_i , B_i , C_i et D_i sont des constantes indépendantes de t , que l'on déterminera grâce aux conditions de continuité suivantes (on a aussi utilisé la condition $\phi(t) \rightarrow 0$ quand $t \rightarrow \pm\infty$) :

$$\begin{cases} \text{en } x'_1 = 0 : \\ A_0 = A_1 \\ \alpha A_1 = \beta B_1 \end{cases} \quad (33)$$

$$\begin{cases} \text{en } y'_{i-1} : \text{ pour } i = 1, \dots, n \\ A_{i-1} \cos(\beta y'_{i-1}) + B_{i-1} \sin(\beta y'_{i-1}) = C_{i-1} e^{\alpha y'_{i-1}} + D_{i-1} e^{-\alpha y'_{i-1}} \\ \beta [-A_{i-1} \sin(\beta y'_{i-1}) + B_{i-1} \cos(\beta y'_{i-1})] = \alpha [C_{i-1} e^{\alpha y'_{i-1}} - D_{i-1} e^{-\alpha y'_{i-1}}] \end{cases} \quad (34)$$

$$\begin{cases} \text{en } x'_i : \text{ pour } i = 2, \dots, n \\ C_{i-1} e^{\alpha x'_i} + D_{i-1} e^{-\alpha x'_i} = A_i \cos(\beta x'_i) + B_i \sin(\beta x'_i) \\ \alpha [C_{i-1} e^{\alpha x'_i} + D_{i-1} e^{-\alpha x'_i}] = \beta [-A_i \sin(\beta x'_i) + B_i \cos(\beta x'_i)] \end{cases} \quad (35)$$

$$\begin{cases} \text{en } y'_n : \\ A_n \cos(\beta y'_n) + B_n \sin(\beta y'_n) = A_{n+1} e^{\alpha y'_n} \\ \beta [-A_n \sin(\beta y'_n) + B_n \cos(\beta y'_n)] = \alpha A_{n+1} e^{\alpha y'_n} \end{cases} \quad (36)$$

0.3.1 Méthode directe

Cette méthode consiste à utiliser toutes les équations sauf la dernière pour calculer, itérativement, tous les A_i , B_i , C_i et D_i en fonction de A_0 (dans les calculs, on posera $A_0=1$, et on normalisera la fonction $\phi(t)$ à la fin des calculs ; la fonction normée sera $\tilde{\phi} = \phi / \|\phi\|$).

Le principe du calcul des coefficients est le suivant : On commence par l'équation (33) pour calculer A_1 et B_1 en fonction de A_0 . Ensuite, l'équation (34) pour $i = 2$ nous donne C_1 et D_1 en fonction de A_1 et B_1 , puis l'équation (35) pour $i = 2$ nous donne A_2 et B_2 en fonction de C_1 et D_1 , etc ... On utilise à la fin l'équation (36) (la première équation seulement) pour obtenir A_{n+1} en fonction de A_n et B_n .

La dernière équation ne sera vérifiée que pour certaines valeurs de E puisque α et β dépendent de E . On obtient ainsi les valeurs propres E_k . Remarquez qu'il y a un nombre

fini de valeurs propres pour les états liés, et que pour chaque valeur propre E_k , on a une seule fonction propre $\phi_k(t)$ (chaque E_k est non dégénérée).

La dernière équation peut s'écrire

$$f(\xi) = \beta [-A_n \sin(\beta y'_n) + B_n \cos(\beta y'_n)] + \alpha A_{n+1} e^{-\alpha y'_n} = 0$$

et pourra être résolue numériquement, avec par exemple, la méthode de la bisection¹³ : On décompose d'abord l'intervalle $\xi \in [0, 1]$ en un grand nombre de petits intervalles $[\xi_m, \xi_{m+1}]$, et on cherche les intervalles qui vérifient $f(\xi_m) f(\xi_{m+1}) < 0$. Ensuite sur chacun de ces intervalles, on applique la méthode de la bisection pour avoir la valeur propre ξ avec une certaine précision que l'on fixera à l'avance (par exemple $\varepsilon = 10^{-15}$).

Une fois que nous avons obtenu une valeur propre ξ , nous calculons la norme de la fonction d'onde correspondante $||\phi||$, en utilisant

$$\begin{aligned} ||\phi||^2 &= \int_{-\infty}^{+\infty} |\phi(t)|^2 dt \\ &= \int_{-\infty}^{x_1} |\phi(t)|^2 dt + \sum_{i=2}^n \int_{y'_{i-1}}^{x'_i} |\phi(t)|^2 dt + \sum_{i=1}^n \int_{x'_i}^{y'_i} |\phi(t)|^2 dt + \int_{y'_n}^{+\infty} |\phi(t)|^2 dt \end{aligned}$$

Cette méthode peut facilement être implémentée en langage C ou bien avec le logiciel **maxima**. Des graphes des niveaux d'énergies ainsi que des fonctions propres correspondantes (convenablement normalisées) peuvent être obtenus.

0.3.2 Méthode de Numerov(differences finies améliorées)

Cette méthode consiste à remplacer l'équation différentielle (32) par une équation aux différences finies, et résoudre algébriquement cette dernière équation. On ré-écrit l'équation (32) comme

$$\phi''(t) - F(t) \phi(t) = 0 \tag{37}$$

où $F(t) \equiv \gamma^2 \left(\xi^2 - U(t) \right)$. Nous ne voulons pas ici calculer $\phi(t)$ pour tous les t mais seulement pour certains t_k vérifiant $t_{k+1} - t_k = h$, h étant une constante choisie suffisamment petite (c'est le pas de discrétisation). Nous divisons l'intervalle des t qui nous intéresse en N petits intervalles $[t_k, t_{k+1}]$, et nous essayons de trouver les $\phi_k \equiv \phi(t_k)$, en commençant par exemple par t_{-1} et t_0 qui sont supposés être choisis par des considérations physiques (c'est une équation différentielle du second ordre, et l'on a besoin de 2 conditions initiales, ou bien de 2 conditions aux limites : tout dépend du problème).

13. La méthode de la bisection est expliquée dans l'Appendice.

L'équation (37) s'écrit au point t_k comme :

$$\phi''(t_k) - F(t_k) \phi(t_k) = 0 \quad (38)$$

et la méthode des différences finies[7][8][9] consiste à remplacer $\phi''(t_k)$ par une approximation qui ne contient plus aucune dérivée. L'approximation désirée est obtenue en utilisant la formule de Taylor[9].

Exemple : si on développe $\phi(t_k - h)$ et $\phi(t_k + h)$ au voisinage du point t_k , on obtient

$$\begin{aligned} \phi(t_k - h) &= \phi(t_k) - h \phi'(t_k) + \frac{h^2}{2!} \phi''(t_k) - \frac{h^3}{3!} \phi'''(t_k) + \mathcal{O}(h^4) \\ \phi(t_k + h) &= \phi(t_k) + h \phi'(t_k) + \frac{h^2}{2!} \phi''(t_k) + \frac{h^3}{3!} \phi'''(t_k) + \mathcal{O}(h^4) \end{aligned}$$

On fait la somme des 2 équations et on en déduit $\phi''(t_k)$

$$\phi''(t_k) = \frac{\phi(t_k + h) - 2\phi(t_k) + \phi(t_k - h)}{h^2} + \mathcal{O}(h^2)$$

En posant $F_k \equiv F(t_k)$, $\phi_k = \phi(t_k)$ et $\phi_{k\pm 1} = \phi(t_k \pm h)$, l'équation (38) devient

$$\phi_{k+1} = (2 + h^2 F_k) \phi_k - \phi_{k-1} + \mathcal{O}(h^4) \quad (39)$$

Cette dernière équation signifie que si nous connaissons ϕ_k et ϕ_{k-1} , on peut calculer ϕ_{k+1} en commettant une erreur locale proportionnelle à h^4 .

La méthode de Numerov¹⁴[10] est une amélioration de cette équation, et elle a une erreur locale proportionnelle à h^6 . La formule¹⁵ ici est :

$$\phi_{k+1} = \frac{\left(2 + \frac{5h^2}{6} F_k\right) \phi_k - \left(1 - \frac{h^2}{12} F_{k-1}\right) \phi_{k-1}}{1 - \frac{h^2}{12} F_{k+1}} + \mathcal{O}(h^6) \quad (40)$$

Pour démarrer, nous devons connaître ϕ_k et ϕ_{k-1} quelque part. Pour notre problème, la fonction $\phi(t)$ est connue pour $t < x'_1$: elle vaut $\phi(t) = A_0 e^{\alpha t}$. En posant A_0 égal à 1 (nous normalisons à la fin des calculs), et comme $x'_1 = 0$, on a :

$$\begin{aligned} \phi_0 &\equiv \phi(0) = 1 \\ \phi_{-1} &\equiv \phi(0 - h) = e^{-\alpha h} \end{aligned} \quad (41)$$

14. Elle est en général utilisée pour résoudre numériquement les équations différentielles de la forme $y''(t) = f(y, t)$ où la dérivée première $y'(t)$ n'apparaît pas. L'équation de Schrödinger est de cette forme avec $f(y, t)$ linéaire en $y(t)$, et cette méthode peut être utilisée en principe pour n'importe quel potentiel.

15. La démonstration de cette méthode est détaillée dans l'Appendice.

Les équations (40) et (41) nous permettent, pour un ξ donné (c.à.d une énergie E donnée), de calculer $\phi_{k+1} \equiv \phi((k+1)h)$ pour $k = 0, 1, \dots, k_{\max}$. $k_{\max}$ est choisi pour que $t_{\max} = (k_{\max} + 1)h \gg y'_n$: Nous voulons dire que nous démarrons en $x'_1 = 0$, intégrons jusqu'à y'_n et allons un peu plus loin pour vérifier si $\phi(t_{\max})$ tend vers 0 ou pas. Si elle tend vers 0, cela voudra dire que le ξ que nous avons pris correspond à un état lié, sinon nous avons une solution qui ne vérifie pas $\phi(t) \rightarrow 0$ quand $t \rightarrow +\infty$. La méthode pratique pour avoir les états liés est donc de poser que $\phi(t_{\max}) = 0$. Cela nous donne une équation à résoudre pour obtenir les valeurs propres.

L'algorithme est comme avant : l'intervalle $[0, 1]$ pour ξ est divisé en un grand nombre de petits intervalles $[\xi_m, \xi_{m+1}]$.

On calcule $\phi_{k_{\max}+1}$ pour ξ_m et ξ_{m+1} . Si on obtient des nombres de signes opposés, alors cet intervalle contient une valeur propre que l'on cherche avec la méthode de la bisection.

0.3.3 Résultats

L'existence des bandes

Avec chacune des deux méthodes nous avons obtenu 26 valeurs propres quasiment identiques :

Directe

-267.289	-266.944	-266.390	-265.666	-264.814	-263.896
-262.986	-262.160	-261.499	-261.073		
-174.531	-173.020	-170.595	167.389	-163.583	-159.406
-155.144	-151.140	-147.803	-145.571		
-36.886	-33.510	-28.028	-20.695	-11.882	-2.212

Numerov

-267.286	-266.941	-266.389	-265.662	-264.809	-263.892
-262.981	-262.157	-261.497	-261.070		
-174.530	-173.018	-170.592	-167.387	-163.581	-159.404
-155.140	-151.136	-147.802	-145.568		
-36.886	-33.509	-28.026	-20.694	-11.880	-2.210

Les énergies obtenues sont arrangées en bandes¹⁶ bien distinctes[11] séparées par des vides qui représentent le gap d'énergie (figure10).

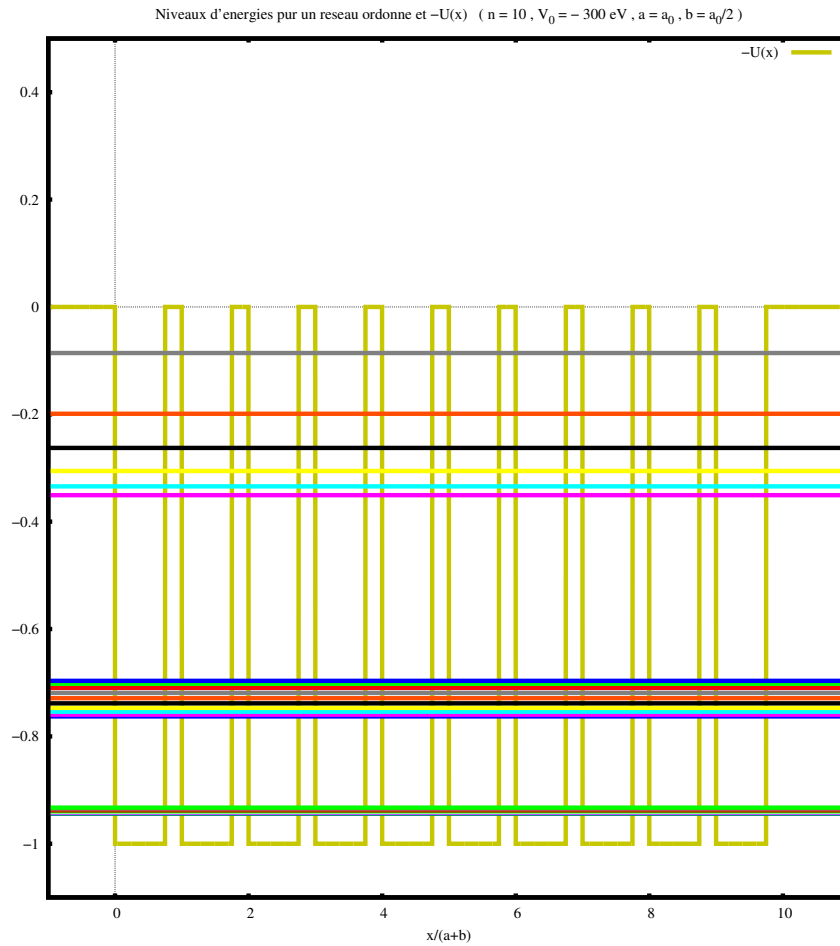


FIGURE 2 – Les niveaux d'énergie pour un réseau régulier de dix puits de potentiel (obtenus avec la méthode directe) : Les niveaux d'énergie sont arrangés en bandes

16. • Dans les figures de ce mémoire, ce sont les valeurs sans dimensions qui sont représentées, c'est-à-dire au lieu des énergies en électron-volt c'est les ξ correspondants qui sont représentés tel que $\xi = \sqrt{E/V_0}$, et au lieu du potentiel $V(x)$ c'est la valeur $U(t)$ tel que $U(t) = V(x)/V_0$.
• Nous avons numéroté les énergies en commençant de zéro : $E_0, E_1, E_2, \dots$

Les fonctions propres

Chaque fonction propre $\phi_m(t)$ s'annule exactement m fois. C'est le théorème des oscillations¹⁷. Nous remarquons que les fonctions d'ondes n'ont pas le caractère d'onde de Bloch supposé au départ par le modèle de Kronig-Penney, mais elles ont la forme d'onde de Bloch stationnaires¹⁸(figure4) c'est à dire le produit d'une onde qui a la périodicité du réseau, et d'une onde stationnaire d'un électron libre. On remarque aussi que pour les fonctions d'ondes des états supérieurs de la dernière bande, on retrouve de plus en plus le caractère périodique des fonctions d'onde, en d'autres termes le caractère de Bloch (figure3). On conclue que les électrons des états supérieurs sont moins liés que ceux des états des premières bandes.

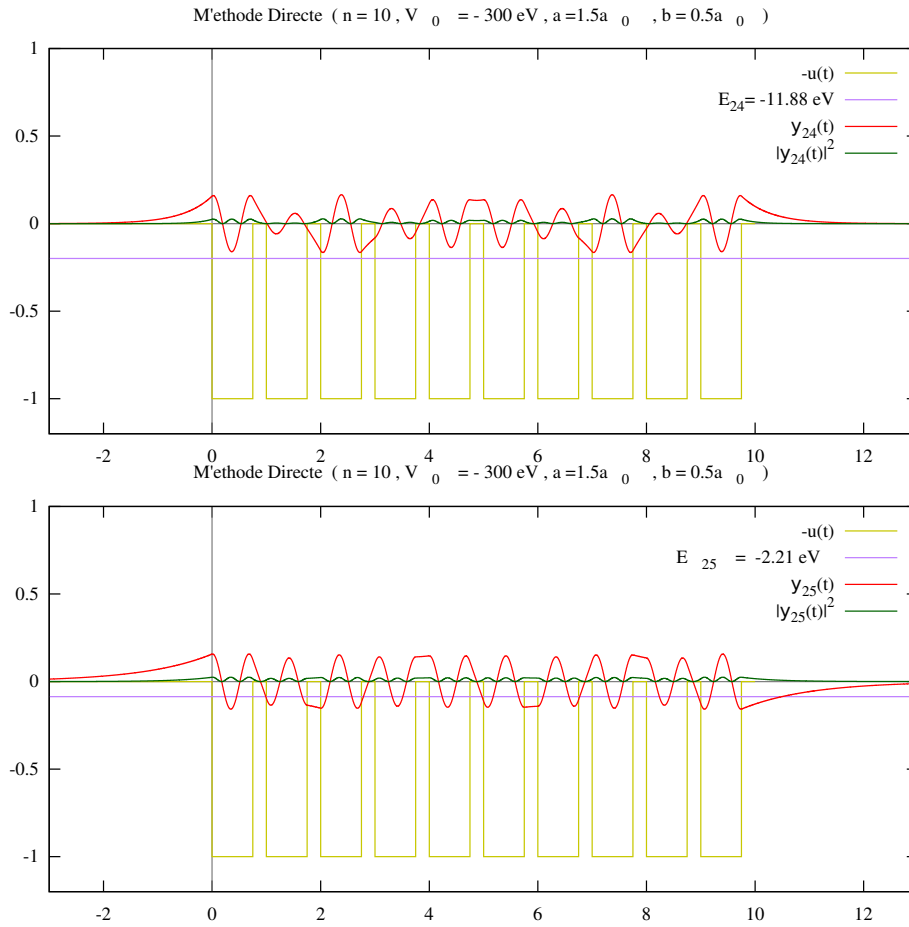


FIGURE 3 – Les fonctions d'ondes des deux derniers états de la dernière bande

17. Le théorème des oscillations[12] stipule que la fonction d'onde $\psi_n(x)$ correspondant à la valeur propre E_n s'annule n fois.

18. Une onde stationnaire est un phénomène résultant de la propagation simultanée dans des directions différentes de plusieurs ondes de même fréquence, dans le même milieu physique. Au lieu d'y avoir une onde qui se propage dans l'espace, on constate une vibration stationnaire mais d'intensité différente en chaque point observé. Les points fixes caractéristiques sont appelés des nœuds de pression

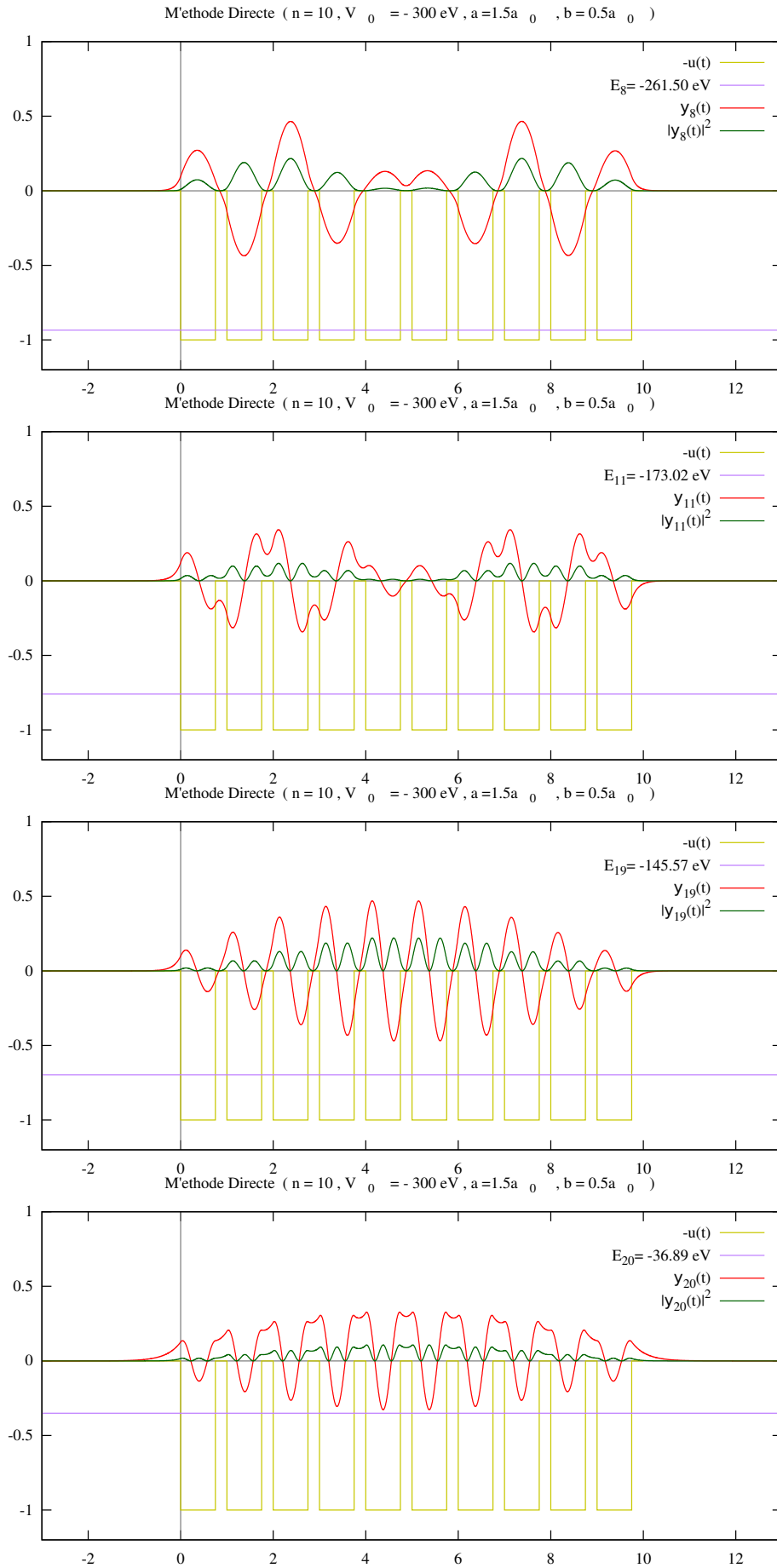


FIGURE 4 – Les fonctions d'ondes et les probabilités de présence de l'électron correspondant à un niveau d'énergie de chaque bande, obtenues avec la méthode directe

Pour la méthode Numerov, on remarque que la fonction d'onde diverge complètement après s'être annulée, et cela est dû au fait que la précision numérique est perdue à la fin du calcul, car à chaque fois que nous calculons la valeur de la fonction d'onde en un point il y a une erreur qui s'accumule, et à la fin du calcul l'erreur devient assez important pour faire diverger la fonction (figure5). Pour limiter le taux d'erreur, il serait plus efficace d'intégrer dans les deux sens (c'est à dire : intégrer de gauche et de droite), de cette façon l'erreur est minimisée et la partie divergente n'apparaît plus.

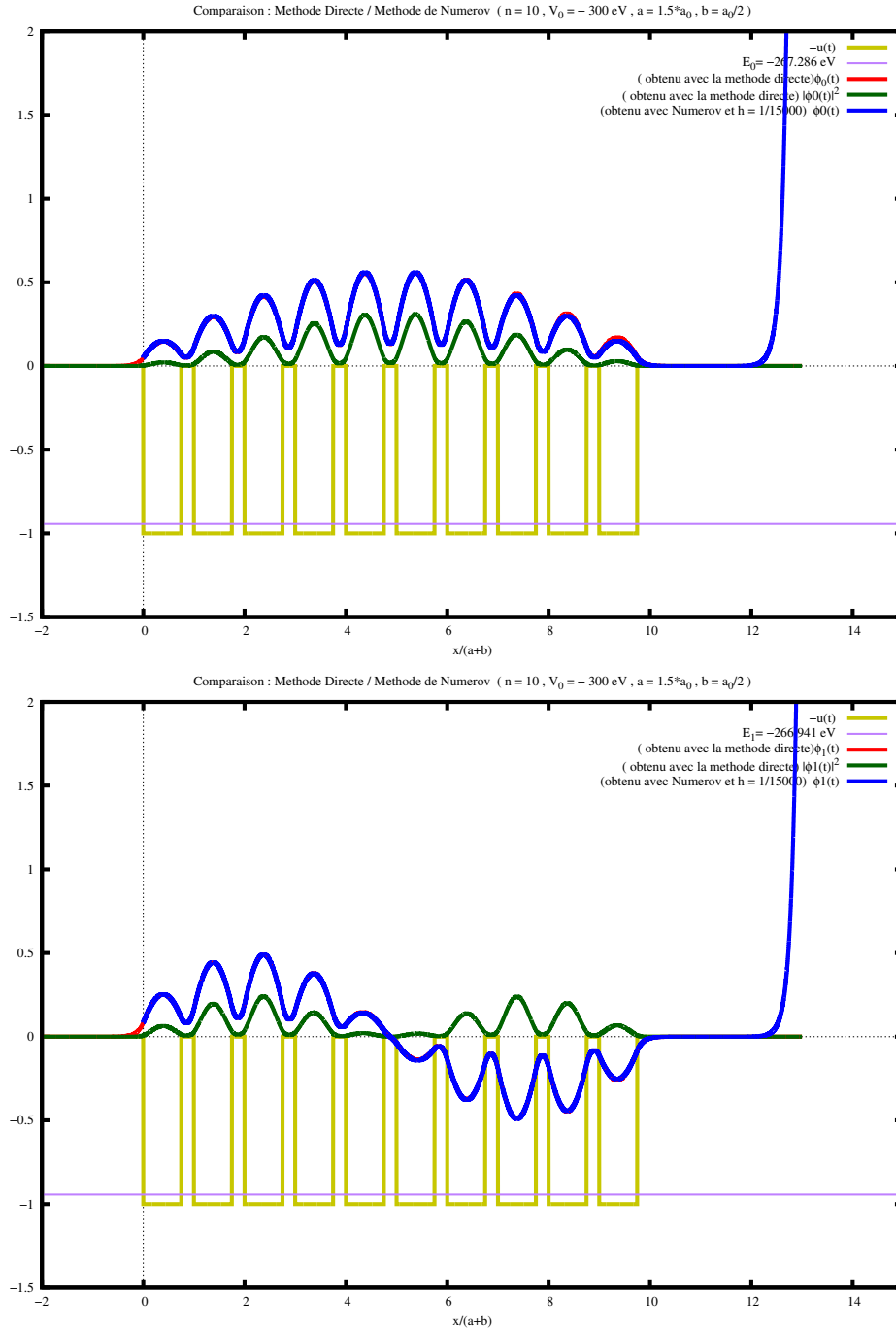


FIGURE 5 – Les fonctions d’ondes et les probabilités de présence de l’électron correspondant à l’état fondamental et les trois premiers états excités des deux méthodes. On voit bien que les graphes obtenus avec les deux méthodes se superposent parfaitement

0.4 Introduction d'une impureté dans le réseau régulier

Maintenant une étape importante dans notre étude est d'introduire une impureté dans le réseau[11]. La façon la plus simple de le faire est de rendre l'un des puits différent des autres (comme une simulation du dopage d'un cristal). Dans notre programme on va diminuer la largeur du cinquième puits d'environ 25%, et on va observer le comportement des fonctions d'ondes et les changements survenus sur les structures de bandes. Pour cette partie on va choisir une seule méthode de calcul, c'est celle de Numerov.

Nous avons obtenu 25 valeurs propres :

-267.038	-266.837	-265.999	-265.310	-264.475	-263.321
-262.798	-261.609	-261.449	-255.620		
-173.448	-172.638	-166.040	-162.118	-157.024	-154.327
-148.533	-147.584	-128.594			
-34.562	-34.562	-24.310	-19.265	-8.473	

On remarque que la structure de bande est maintenue, cependant, les niveaux d'énergie sont légèrement translatés, le dernier état de la dernière bande est devenu un état non lié d'où les 25 valeurs propres au lieu des 26 trouvées dans le réseau sans impuretés. On remarque aussi qu'il y a des niveaux d'énergie qui se retrouvent dans la bande interdite (figure 6), ce sont les états 10 et 20 d'énergie $-255,62$ eV et $-128,59$ eV respectivement. Ces états à l'intérieur du gap, sont dans un sens indépendants (ou découplés) du reste du réseau. Les états propres du puits modifié se positionnent à l'endroit où ils seraient si le reste du réseau était inexistant, d'une autre manière se sera comme si le réseau avait juste un seul puit. Donc le positionnement de ces états dépend juste de l'endroit où se situe l'impureté.

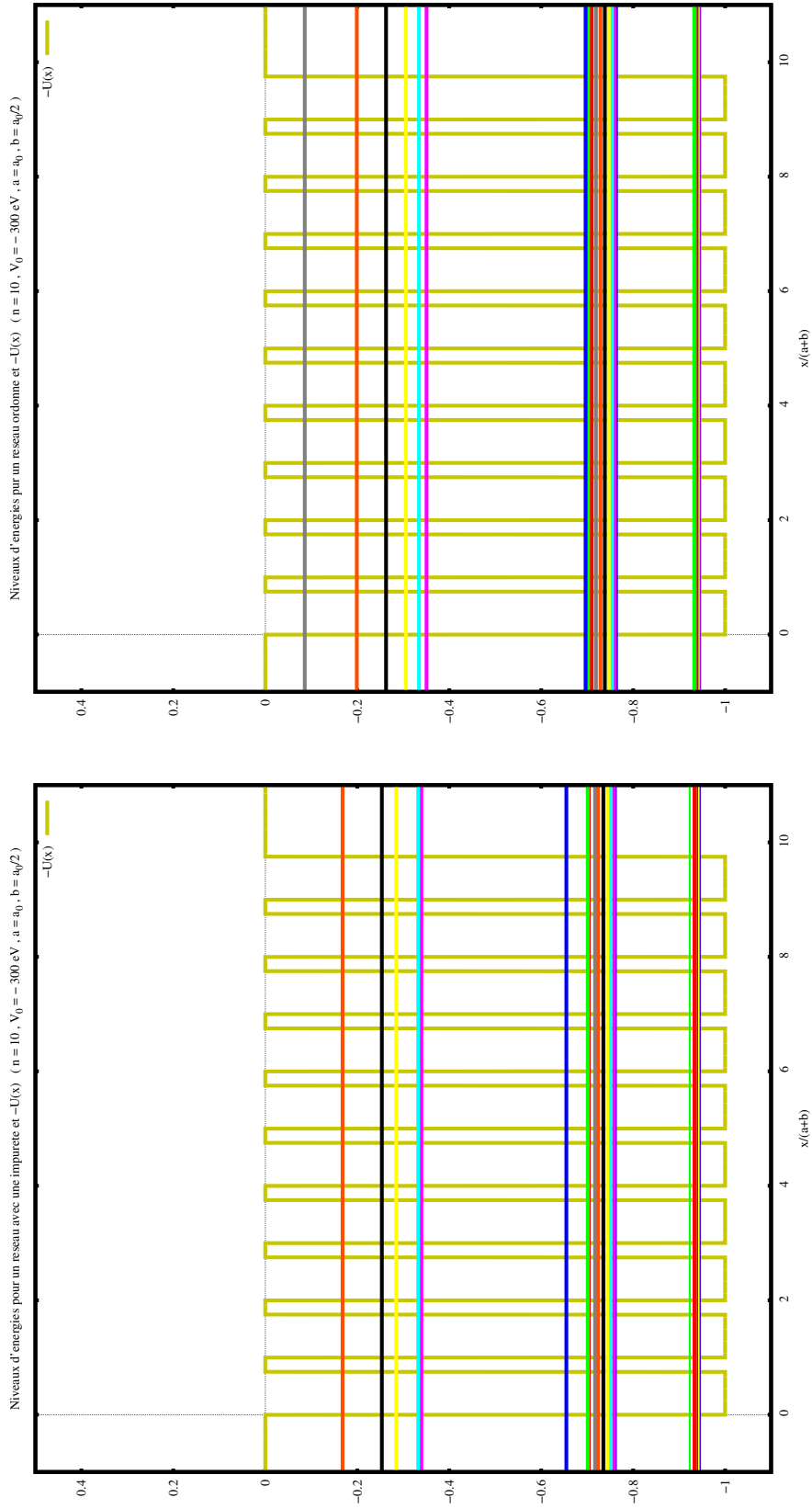


FIGURE 6 – A droite, les niveaux d'énergie du réseau régulier. A gauche, le réseau avec impuretés : On voit bien l'apparition de niveau d'énergie à l'intérieur du gap

La localisation de la fonction d'onde

Si nous examinons les formes des fonctions d'ondes, on trouve que la présence de cette impureté a détruit la structure d'onde de Bloch des fonctions d'ondes. Certaines de ces fonctions se propagent le long du réseau, d'autres sont positionnées à gauche ou à droite de l'impureté (figure7), mais les fonctions correspondants aux niveaux translatés à l'intérieur des gaps sont localisés autour de l'impureté et gardent la même forme pour n'importe quelle position de l'impureté à l'intérieur du réseau (figure8).

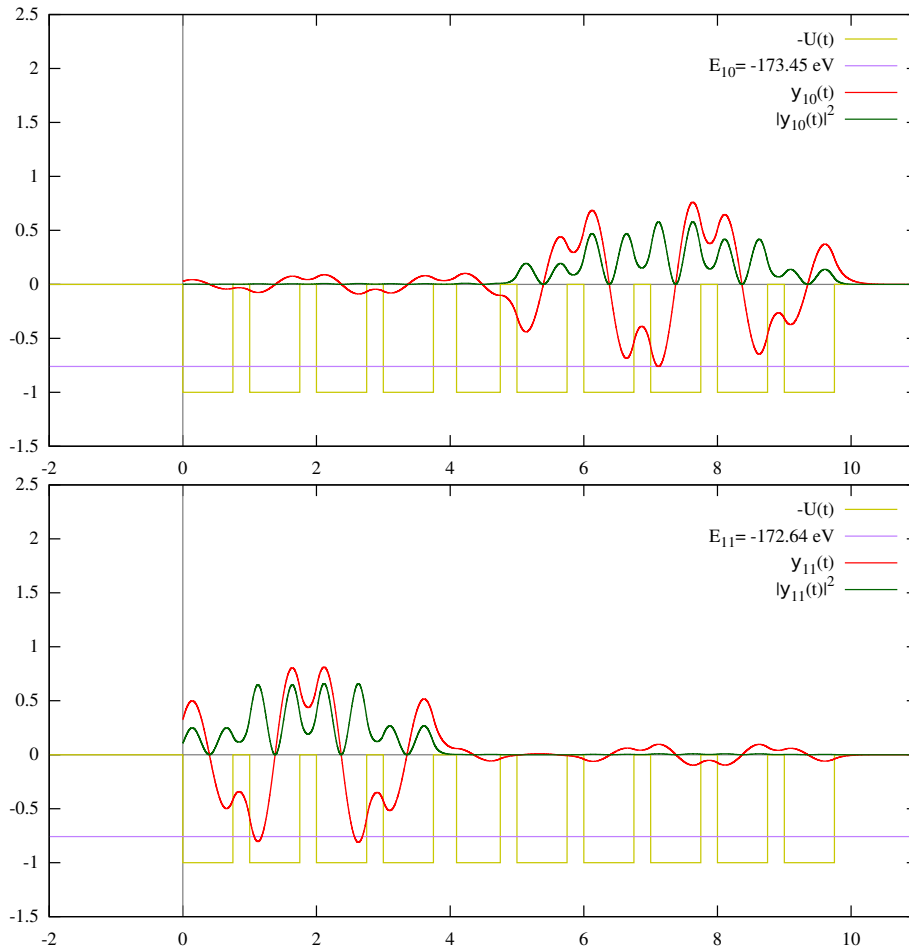


FIGURE 7 – Les fonctions d'ondes des niveaux situés à l'intérieur des bandes : On voit bien que les fonctions sont situées à gauche ou à droite de l'impureté

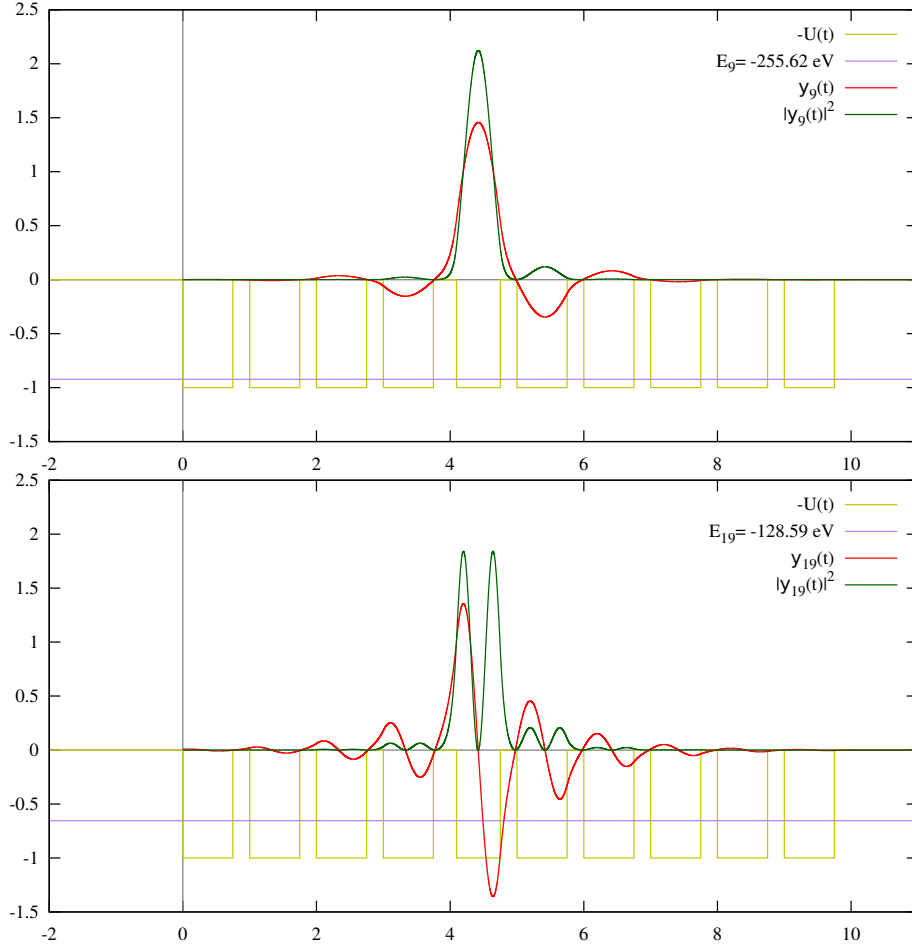


FIGURE 8 – Les fonctions d’ondes et probabilités de présence des états situés dans les gaps : Les fonctions d’ondes sont localisées autour de l’impureté

Comme mentionné dans la partie théorique le dopage a un grand intérêt pour l’amélioration de la conductivité électrique dans les semiconducteurs, et dans la présente étude, même pour un système quelconque comme le notre, nous avons démontré qu’en introduisant des impuretés dans un solide des états apparaissaient dans les gaps précédemment inoccupés, en d’autres termes c’est comme si nous avions créé des niveaux donneurs ou accepteurs à l’intérieur du gap.

0.5 Réseau désordonné

Dans le modèle original d’Anderson, le désordre dans le réseau cristallin se traduisait dans la variation de la profondeur des puits, en d’autres termes une variation du potentiel de chaque atome du réseau. Dans notre modélisation du problème, nous avons choisi de simuler le désordre en changeant la distance entre les puits[11] pour qu’elle soit différente entre chaque deux puits. Pour y arriver numériquement, nous avons utilisé un générateur

de nombres aléatoires.

Pour le potentiel généré, nous avons obtenu 26 valeurs d'énergie :

-270.99	-266.31	-265.50	-265.12	-264.07	-262.58
-260.25	-258.25	-181.88	-180.88	-172.76	-166.61
-163.51	-161.97	-156.10	-152.79	-150.65	-148.36
-137.11					
-44.34	-39.18	-34.61	-25.89	-14.02	-4.85
-4.38					

La figure 9 illustre la forme du potentiel généré ainsi que les 26 valeurs d'énergie trouvées réparties en bandes.

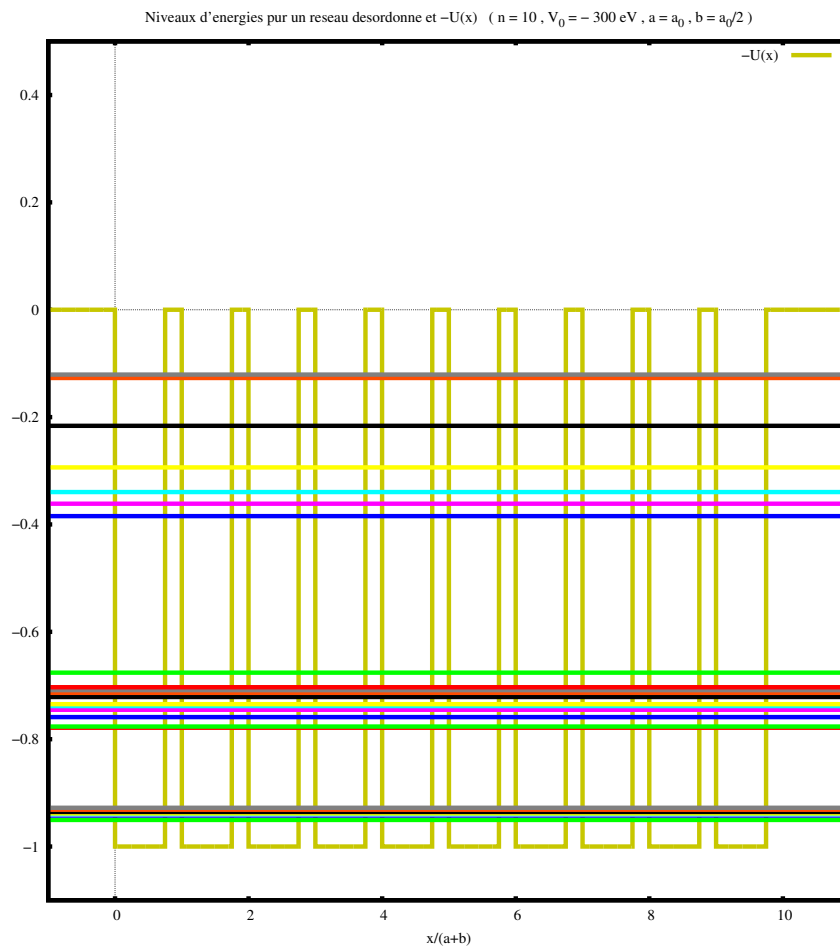


FIGURE 9 – Les niveaux d'énergie pour un réseau irrégulier de dix puits de potentiel

Le nombre exacte de niveaux d'énergie, et leur répartition en structure de bande peuvent varier selon la disposition des puits de potentiel, néanmoins la physique reste la même.

Nous avons constaté que l'énergie a conservé son caractère de bande, cependant, les niveaux d'énergie ont été déplacés de sorte que certains de ces niveaux se sont retrouvés dans les gaps. Nous avons aussi constaté que le caractère d'onde de Bloch des fonctions d'ondes est largement perturbé par le désordre, on ne voit plus des états étendus sur le réseau, mais des états localisés autour des puits (figure10), ce qui confirme la théorie d'Anderson d'états localisés dans les solides amorphes.

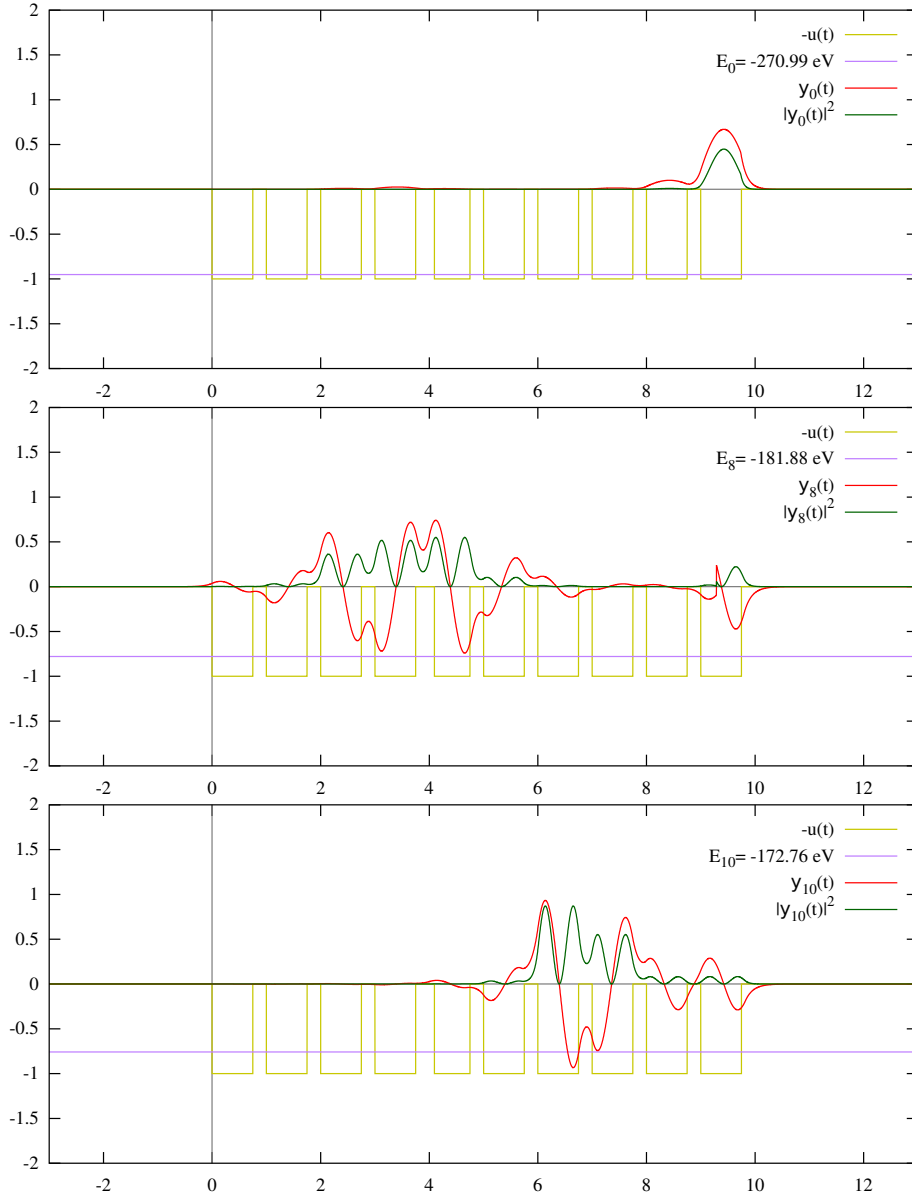


FIGURE 10 – Les fonctions d’ondes et les probabilités de présence de l’électron correspondant aux premiers niveaux des trois premières bandes d’énergie. On voit que les fonctions d’onde sont localisées

En analysant les fonctions d’onde dans chaque bande, nous avons réalisé que les états correspondants aux bandes d’énergie les plus basses ont plus tendance à être localisés que les états des plus hauts niveaux, qui quand à eux, ont une partie de la fonction d’onde qui est étendue et garde son caractère d’onde de Bloch comme dans un réseau régulier (figure11).

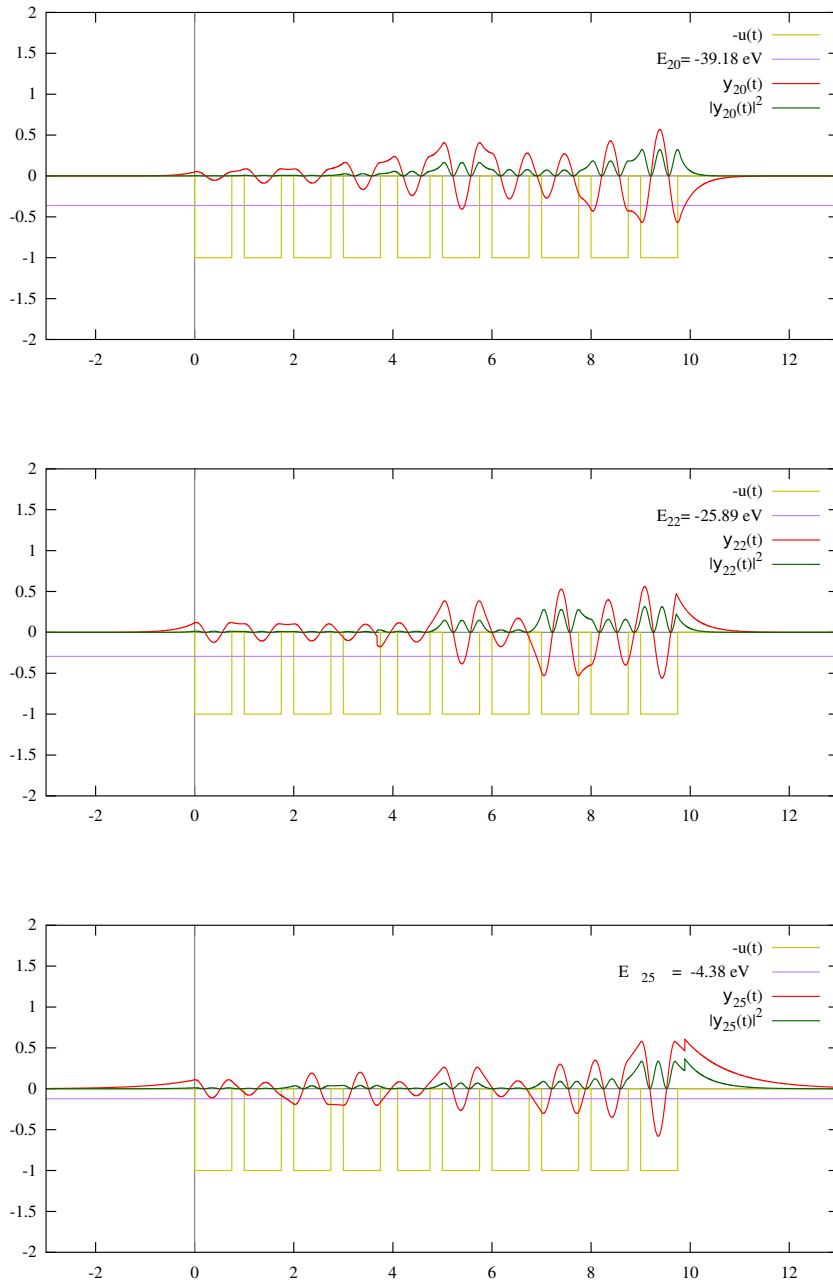


FIGURE 11 – Les fonctions d’ondes et les probabilités de présence de l’électron correspondant aux niveaux d’énergies des dernières bandes. On voit que les fonctions d’ondes sont délocalisées comme dans le cas d’un réseau régulier

Nous avons aussi constaté, que les états se trouvant vers les bords des bandes sont plus localisés que ceux vers le centre (figure12).

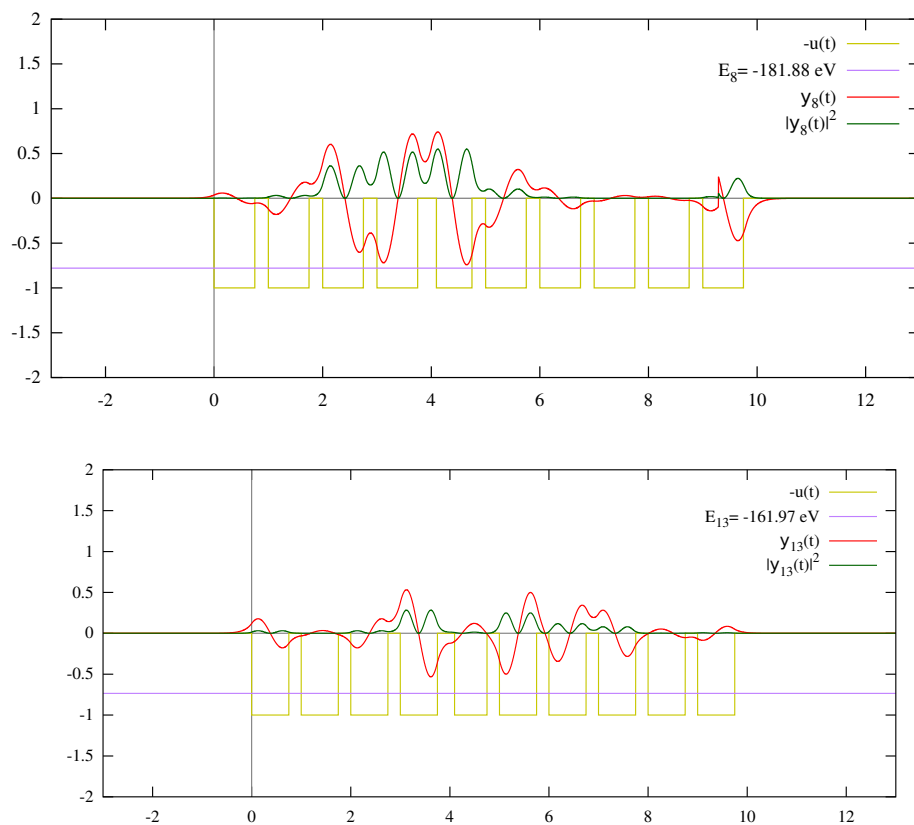


FIGURE 12 – En haut, on a un état localisé du bord de bande (le premier niveau de la deuxième bande). En bas parcontre, on voit un état étendu du centre de la bande(un état du centre de la deuxième bande).

Conclusion

Dans ce mémoire, nous avons étudié les propriétés d'un électron dans un réseau cristallin à une dimension. Ainsi, pour un réseau périodique de puits carrés de potentiel, nous avons constaté que les énergies propres de cet électron constituaient des bandes d'énergies et les fonctions d'onde associées avaient la forme d'onde stationnaires qui s'étendaient sur tout le réseau.

Nous avons modifié l'un des puits pour simuler l'effet d'un dopage ; nous avons remarqué que des états sont apparus dans le gap, et ces états là étaient fortement localisés autour de l'impureté.

Pour un réseau complètement désordonné, nous avons obtenu des états complètement localisés, vérifiant ainsi l'hypothèse d'Anderson. Dans un vrai matériau cet effet de localisation causé par le désordre se traduirait par une grande diminution de la conductivité rendant ainsi le matériau isolant.

La théorie des bandes a permis d'expliquer beaucoup de propriétés physiques et chimiques des matériaux. Ainsi par le taux de remplissage des bandes on peut déterminer la nature du matériau en question : Conducteur, semiconducteur ou isolant. Cependant, c'est une théorie à un seul électron, et elle n'est plus adéquate pour décrire des systèmes en interactions et où on doit prendre compte des corrélations électroniques, et surtout elle atteint ces limites quand il s'agit d'expliquer certains phénomènes comme la supraconductivité, et le cas des matériaux isolants considérés comme des métaux par la théorie des bande (isolant de Mott). Pour ces situations d'autres méthodes sont plus adaptées telle que la méthode de Monte-carlo quantique ou la DMFT (Dynamical Mean Field Theory) qui est plus efficace pour déterminer la structure électronique des matériaux fortement corrélés.

Appendices

.1 La méthode de Dichotomie ou Bissection

La dichotomie ou communément appelée bissection[9], est l'une des méthodes les plus simples pour trouver le zéro d'une fonction. Le principe de cette méthode est le suivant : nous voulons trouver le zéro d'une fonction f dans l'intervalle $[a, b]$ (la fonction f est continue sur cet intervalle et $f(a)$ et $f(b)$ sont de signes opposés), en d'autres termes nous allons chercher les x pour lesquels $f(x) = 0$. On va procéder comme suit :

On divise l'intervalle $[a, b]$ en deux sous-intervalles, et on calcule $c = (a + b)/2$. Il y a maintenant trois possibilités :

- **la première** : $f(c) = 0$ (... et c'est comme gagner à la loterie), on dit que c est une solution exacte, et là le travail est terminé !
- **la deuxième possibilité** : $f(a)f(c) < 0$; le zéro se trouve dans l'intervalle $[a, c]$.
- **la troisième et la dernière** : $f(c)f(b) < 0$; le zéro se trouve dans l'intervalle $[c, b]$.

Quand on a isolé le sous intervalle où se trouve la solution, on refait les mêmes étapes que précédemment pour le sous-intervalle en question, et cela jusqu'à obtention de la précision désirée (une précision fixée à l'avance).

Pour plus de clarté voici un algorithme qui résume approximativement ce qui a été dit :

On a une fonction f continue sur l'intervalle $[a, b]$, $c = (a + b)/2$:

si $f(c) = 0 \implies c$ est une solution exacte, le travail est terminé.

Si non :

si $f(a)f(c) < 0 \implies a \leftarrow a$ et $b \leftarrow c$

si $f(a)f(c) > 0$ (d'une autre manière : si $f(b).f(c) < 0$) $\implies a \leftarrow c$ et $b \leftarrow b$

Cette procédure est répétée jusqu'à ce que : $|(b - a)/2| < \varepsilon$, avec ε une précision fixée à l'avance (par exemple : $\varepsilon = 10^{-15}$).

De cette façon, on constate que la bissection est une méthode récursive, et que la solution n'est pas obtenue directement par un seul calcul.

.2 La méthode de Numerov

La méthode Numerov[13] est une méthode numérique pour résoudre les équations différentielles ordinaires du second ordre, dans lesquelles le terme du premier ordre n'apparaît pas ; il s'agit d'une méthode multi-étapes (itérative) du 4^{ème} ordre. La procédure de calcul commence par un point initial et prend un pas pour trouver le point suivant, le processus se poursuit pour les autres points, jusqu'à ce qu'on trace la fonction. Cette méthode est implicite(ie : donne des solutions approchées), mais elle peut être explicite si l'équation différentielle est linéaire. La méthode Numerov était développée par l'astronome russe « Boris Vasilyevich Numerov ».

Le principe de la méthode :

La méthode Numerov est utilisée pour résoudre des équations différentielles de la forme :

$$\left(\frac{d^2}{dx^2} + f(x) \right) y(x) = 0 \quad (42)$$

la fonction $y(x)$ est calculée en tous points équidistant x_n de l'intervalle $[a, \dots, b]$ tels que $x_n - x_{n-1} = h$; h est le pas ou la distance entre deux points consécutifs. Pour commencer, on doit avoir deux valeurs initiales de la fonction $y(x)$ en deux points consécutifs, et le reste des valeurs de la fonction est calculée par :

$$y_{n+1} = \frac{\left(2 - \frac{5h^2}{6} f_n \right) y_n - \left(1 + \frac{h^2}{12} f_{n-1} \right) y_{n-1}}{1 + \frac{h^2}{12} f_{n+1}} \quad (43)$$

tel que : $f_n = f(x_n)$ et $y_n = y(x_n)$ sont les valeurs des fonctions au point x_n .

Pour arriver au résultat précédent, on utilise la formule de Taylor :

si on développe $y(x_k - h)$ et $y(x_k + h)$ au voisinage du point x_k , on obtient :

$$\begin{aligned} y_{k-1} &\equiv y_{x_k-h} \\ &= y(x_k) - hy'(x_k) + \frac{h^2}{2!}y''(x_k) \\ &\quad - \frac{h^3}{3!}y'''(x_k) + \frac{h^4}{4!}y^{(4)}(x_k) - \frac{h^5}{5!}y^{(5)}(x_k) + \mathcal{O}(h^6) \end{aligned}$$

$$\begin{aligned} y_{k+1} &\equiv y_{x_k+h} \\ &= y(x_k) + hy'(x_k) + \frac{h^2}{2!}y''(x_k) \\ &\quad + \frac{h^3}{3!}y'''(x_k) + \frac{h^4}{4!}y^{(4)}(x_k) + \frac{h^5}{5!}y^{(5)}(x_k) + \mathcal{O}(h^6) \end{aligned}$$

On additionne les deux expressions :

$$y_{k+1} + y_{k-1} = 2y_k + h^2y''_k + \frac{h^4}{12}y^{(4)}_k + \mathcal{O}(h^6) \quad (44)$$

On ramène y''_k au premier membre :

$$-h^2y''_k = -y_{k-1} - y_{k+1} + 2y_k + \frac{h^4}{12}y^{(4)}_k + \mathcal{O}(h^6) \quad (45)$$

On remplace y''_k par l'expression $-f_k y_k$ déduite de notre équation différentielle (42) de départ, on obtient :

$$h^2 f_k y_k = -y_{k-1} + 2y_k - y_{k+1} + \frac{h^4}{12}y^{(4)}_k + \mathcal{O}(h^6) \quad (46)$$

Maintenant on prend la dérivé seconde de notre équation différentielle (42) :

$$\frac{d^2}{dx^2} \left(\frac{d^2}{dx^2} + f(x) \right) y(x) = 0 \quad (47)$$

$$\begin{aligned} \frac{d^2}{dx^2} \left(\frac{d^2}{dx^2} y(x) \right) &= -\frac{d^2}{dx^2} \left(f(x)y(x) \right) \\ y^{(4)}(x) &= -\frac{d^2}{dx^2} \left(f(x)y(x) \right) \end{aligned} \quad (48)$$

On remplace la dérivé seconde, du second membre de l'équation (48) par son quotient

différentiel du second ordre, et on l'injecte dans l'équation (46) :

$$h^2 f_k y_k = -y_{k-1} + 2y_k - y_{k+1} - \frac{h^4}{12} \left(\frac{f_{k-1} y_{k-1} - 2f_k y_k + f_{k+1} y_{k+1}}{h^2} \right) + \mathcal{O}(h^6) \quad (49)$$

En ramenant y_{k+1} au premier membre, et en prenant le y_k et le y_{k-1} en facteur au deuxième membre, on aura :

$$y_{k+1} = \frac{\left(2 - \frac{5h^2}{6} f_k\right) y_k - \left(1 + \frac{h^2}{12} f_{k-1}\right) y_{k-1}}{1 + \frac{h^2}{12} f_{k+1}} \quad (50)$$

Et on est arrivé à l'équation (43) !

Bibliographie

- [1] Von Felix Bloch. Über quantenmechanik der elektronen in kristallgittern. *Z.Physik*, 1928.
- [2] R. de L.Kronig and W.G. Penney. Quantum mechanics of electrons in crystal lattices. *Proc. R. Soc. Lond.*, 1930.
- [3] P.W.Anderson. Absence of diffusion in certain random lattices. *Physical Review*, 1958.
- [4] Fung Duan and Jin Guo Jun. *Introduction to the condensed matter physics*. World Scientific Publishing Co. Pte. Ltd., March. 2005.
- [5] Claude Aslangul. *Application de la mécanique quantique : De l'atome au solide*. Université Pierre et Marie Curie (Paris 6), 2005/2006.
- [6] Sheng S.Li. *Semiconductor physical electronics*. Springer Science and Business Media, LLC, 2006.
- [7] Joe D.Hoffman. *Numerical Methods for Engineers and Scientists, 2nd Ed.* Marcel Dekker, Inc.,New York, 2001.
- [8] Zhilin Li. *Finite Differences Methods Basics*. Center for Research in Scientific Computation & Department of Mathematics North Carolina State University.
- [9] Adel Kassa. Cours de méthodes numériques, 13 novembre 2008.
- [10] Philippe De Forcrand and Philipp Werner. *Computational Quantum Physics*. ETH Zürich.
- [11] I.D.Jonston and D.Segal. Electron in a crystal lattice : A simple computer model. *Am.J.Phys.*, Vol.60, No.7, July 1992.
- [12] L.D Landau and E.M Lifshitz. *Quantum Mechanic, Non-Relativistic Theory*. Pergamon Press.
- [13] [http ://en.wikipedia.org/wiki/Numerov's method](http://en.wikipedia.org/wiki/Numerov's_method).