

Nanomonde, nanotechnologies

J.P. Nougier

Professeur émérite à l'Université Montpellier 2, Institut d'Electronique du Sud (UMR CNRS 5214)

Conférence donnée le 19 Novembre 2007 à l'Académie des Sciences et Lettres de Montpellier

	Page
1. Nanomètre, nanomonde, nanotechnologies	2
2. Hier et aujourd'hui : de la micro à la nanoélectronique	2
2.1. En route vers la nanoélectronique	2
2.2. La microélectronique à l'assaut du monde	4
3. Généralités sur la mécanique quantique	6
4. La dualité onde-corpuscule : du microscope optique au microscope électronique	7
4.1. Dualité onde-corpuscule	7
4.2. Microscopes électroniques	9
5. L'effet tunnel et ses applications	10
5.1. Description de l'effet tunnel	10
5.2. Microscope à effet tunnel	11
5.3. Microscope à Force Atomique (AFM)	13
6. Conclusion	14

Des informations très intéressantes sont exposés sous une forme très pédagogique sur le site <http://www.nanomicro.recherche.gouv.fr> de la Direction de la Recherche du Ministère de l'Enseignement Supérieur et de la Recherche.

En particulier, il est vivement conseillé de télécharger la remarquable plaquette de 32 pages intitulée "A la découverte du nanomonde" sur le site :

<http://www.nanomicro.recherche.gouv.fr/nanomonde/nanoplaquette.pdf>.

Cette plaquette est complétée par des animations qui l'illustrent de façon tout à fait attractive et vivante sur le site <http://www.nanomicro.recherche.gouv.fr/nanomonde/nanomonde.swf>. Ce site illustre ainsi l'intérêt que l'on peut tirer d'une pédagogie utilisant les technologies de l'information et de la communication.

Compte tenu de l'existence de ces documents, j'ai décidé d'une part de me limiter à quelques aspects liés à la physique, en éliminant les notions pourtant essentielles impliquant la chimie et la biologie, que je connais moins bien, d'autre part de donner à la présente conférence un aspect totalement différent et complémentaire du contenu des documents référencés ci-dessus, en essayant notamment d'y introduire de façon aussi simple et imagée que possible des notions de physique réputées complexes qui sont à la base des phénomènes intervenant dans les structures de dimensions nanométriques.

"Si vous ne pouvez expliquer un concept à un enfant de six ans, c'est que vous ne le comprenez pas complètement" (Albert Einstein).

1. Nanomètre, nanomonde, nanotechnologies

Un nanomètre (nm) est égal à un milliardième de mètre. C'est une dimension si petite qu'on a du mal à l'imaginer. Elle représente un millionième de millimètre, ou encore un millième de micron. On manipule couramment des milliards d'euros en économie sans se rendre compte de ce que cela représente. A titre d'exemples : entre Lille et Perpignan, distants d'environ 1000 km, il y a un milliard de millimètres ; un milliard de secondes s'écoulent en 31 années ; en l'an 2000, un milliard de minutes s'étaient écoulées depuis la naissance du Christ. Cependant le nanomètre est dans la nature une dimension tout à fait courante, qui représente la taille de quelques atomes, la dimension de petites molécules.

En ce sens, nous pouvons dire que nous baignons dans un "nanomonde" tout aussi naturellement que monsieur Jourdain faisait de la prose puisque, pour parler simplement, le monde est constitué d'atomes.

Un organisme vivant, une cellule, est donc composée d'éléments de dimensions nanométriques qui emmagasinent et dépensent de l'énergie, se déplacent, transportent des objets, organisent des échanges, font vivre la cellule. Ce sont des moteurs moléculaires : moteurs linéaires qui se déplacent le long de filaments rigides ; moteurs rotatifs, qui tournent dans une membrane cellulaire ; moteurs engendrant des mouvements oscillatoires, permettant la nage de certains organismes unicellulaires. Enfin, il y a des molécules qui se déplacent le long de la double hélice de l'ADN, le porteur du code génétique : ces molécules ouvrent l'hélice, dupliquent le code ou créent une copie sur un brin d'ARN. Des techniques de micromanipulation permettent d'étudier leur fonctionnement. Mieux, il est maintenant possible de créer artificiellement des systèmes moléculaires actifs qui représentent un premier pas vers une technologie de moteurs moléculaires, ayant des dimensions de quelques nanomètres.

On appelle nanotechnologies l'ensemble des techniques permettant de créer artificiellement des structures dont l'une au moins des trois dimensions est de l'ordre de quelques nanomètres. En fait on commence à parler de structures nanométriques (et donc de nanotechnologies) lorsqu'une dimension au moins est plus petite qu'une centaine de nanomètres soit $0,1 \mu\text{m}$ (μm se lit "micromètre" ou "micron"). Ces techniques concernent les domaines de la biologie, de la chimie, de la physique. Je me limiterai dans cette conférence à quelques exemples choisis dans le domaine de la physique.

2. Hier et aujourd'hui : de la micro à la nanoélectronique

2.1. En route vers la nanoélectronique

L'ancêtre des microtechnologies est indiscutablement constitué par le transistor, inventé en 1948 par John Bardeen, Walter Brattain et William Shockley (laboratoires Bell, USA) qui leur a valu le prix Nobel en 1956.

Le transistor est un composant électronique qui possède trois électrodes : source, grille, drain (figure F2.1a) : les électrons circulent de la source vers le drain, le courant qu'ils génèrent est contrôlé par la grille. On peut comparer son comportement à celui d'un tuyau où l'eau circule entre un réservoir (la source) et l'embouchure (le drain), le courant d'eau est contrôlé par un robinet ou une vanne (la grille).

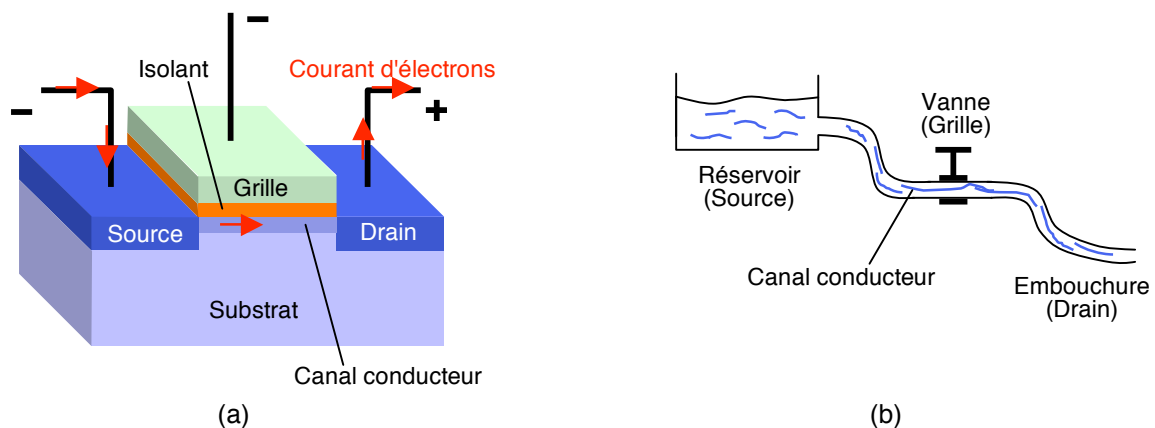


Figure F2.1 : Schéma d'un transistor (a) et analogie hydraulique (b)

– Si l'on n'applique aucune tension sur la grille, le transistor se comporte comme le tuyau dont la vanne est ouverte : si l'on applique une tension positive sur le drain, les électrons sont attirés vers le drain et circulent

de la source vers le drain ; si le réservoir est situé au dessus de l'embouchure, l'eau coule dans le tuyau du réservoir vers l'orifice. Plus la tension positive appliquée sur le drain est forte, plus le courant d'électrons est important ; plus on élève le réservoir par rapport à l'embouchure, plus le courant d'eau est fort dans le tuyau. Si l'on applique sur le drain une tension négative, elle repousse les électrons, qui ne peuvent plus circuler : si l'on abaisse le réservoir au dessous du niveau de l'orifice, l'eau ne circule plus dans le tuyau.

– Si l'on polarise la grille négativement, une partie des électrons est repoussée et n'atteint pas le drain, on diminue le courant électrique, si la polarisation de grille est très négative, tous les électrons sont repoussés, le circuit est coupé. C'est exactement ce qui se passe dans le tuyau si l'on ferme plus ou moins la vanne.

Du fonctionnement qui vient d'être décrit résultent les deux fonctions fondamentales du transistor :

– Fonction amplificatrice : si l'on applique sur la grille une tension sinusoïdale, les électrons venant de la source sont plus ou moins repoussés, on obtient donc un courant drain sinusoïdal, dont l'amplitude est beaucoup plus grande que l'amplitude de la sinusoïde appliquée sur la grille. C'est exactement le même phénomène qui se produit lorsqu'un opérateur tourne alternativement dans un sens puis dans l'autre la vanne d'une immense conduite d'eau qui alimente une centrale électrique : le débit du courant dans le tuyau varie lui aussi en fonction de la position de la vanne : un débit d'eau très important, donc développant une puissance considérable, peut être ainsi contrôlé par la faible puissance développée par l'opérateur actionnant la vanne : la variation de puissance de l'opérateur est donc amplifiée à l'embouchure du tuyau. Ainsi le transistor est utilisé comme amplificateur dans un très grand nombre de circuits analogiques (chaînes haute fidélité, etc.) ;

– Fonction logique : si la tension appliquée sur la grille est suffisamment négative, le courant drain est coupé, on a un fonctionnement en "tout ou rien", c'est un circuit logique, un interrupteur ouvert ou fermé suivant la commande appliquée sur la grille. C'est cette fonction logique qui est utilisée dans tous les circuits numériques, les ordinateurs, etc.

Les circuits électroniques consistent en assemblages de composants, en particulier transistors, résistances, capacités, etc. interconnectés entre eux. L'interconnexion des premiers circuits était réalisée par des fils. On utilisa ensuite des circuits imprimés constitués de plaquettes de bakélite où sont tracées (imprimées) des pistes en cuivre.

Le circuit intégré a été inventé en 1959 par Jack Kilby de Texas Instruments, prix Nobel de physique en 2000 et Robert Noyce (Fairchild Semiconductor) : ici tous les composants sont réalisés sur un même support ainsi que les interconnexions. Aujourd'hui on extrait du sable le Silicium qu'on purifie et qu'on "tire" en lingots de 30 cm de diamètre. Le lingot est coupé en tranches (wafers) d'environ 1 mm d'épaisseur. En utilisant des procédés de photolithographie (analogues au tirage des photos argentiques) et différents traitements, on construit sur une tranche un grand nombre de "puces" d'environ 1 cm² (figure F2.2), chacune d'elles constitue un processeur ou un microcontrôleur constitué de milliers de transistors. C'est une telle puce que l'on voit sur une carte bancaire, une carte vitale, la carte mémoire d'un appareil photo numérique, etc. Actuellement une puce peut comporter 100 millions de transistors, chaque transistor et ses connexions tient donc dans un carré de 10 μ m de côté, de sorte que la largeur de la grille d'un transistor peut être largement inférieure au micron, de l'ordre de quelques dizaines de nanomètres (figure F2.3).

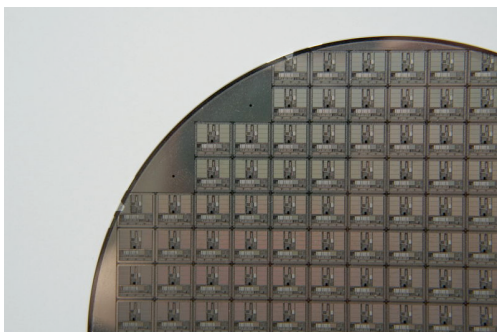


Figure F2.2 : Morceau de tranche de silicium sur laquelle ont été fabriqués des microprocesseurs (Source : Wikipedia n° 0214)

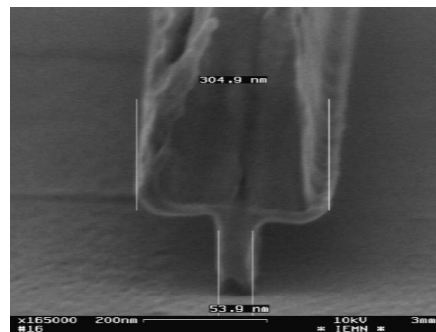


Figure F2.3 : Photographie de la grille d'un transistor, de largeur à la base 54 nm (Source : S. Bollaert, IEMN, Univ. Lille 1, HEMTs et composants AlInAs/GaInAs, AS CNRS n° 207 "Sources et Systèmes THz", Sept. 2004)

Pour obtenir de telles performances, il est nécessaires de procéder aux différentes étapes (près d'une centaine) de fabrication des circuits intégrés dans des salles blanches à l'abri de toute poussière (figure F2.4).

Ainsi l'aventure fabuleuse de la microélectronique consiste depuis 50 ans à réduire de plus en plus les dimensions, pour atteindre aujourd'hui la nanoélectronique. Gordon Moore, co-fondateur de la société Intel (USA), avait prédit que le nombre de transistors par puce doublerait tous les 18 mois, que dans le même laps de temps les dimensions des grilles des transistors serait réduit d'un facteur 1,3 : cette "loi" ne s'est pas démentie depuis 30 ans.

Conséquence de cette réduction des dimensions : le premier ordinateur construit avec des tubes à vides, pendant la seconde guerre mondiale pour modéliser la bombe atomique, occupait 35 m², coûtait plusieurs milliards de dollars, dégageait une énorme chaleur, tombait constamment en panne : les calculatrices scientifiques que l'on achète aujourd'hui dans n'importe quel supermarché pour quelques dizaines d'euros pèsent quelques dizaines de grammes, sont beaucoup plus puissantes que ce premier ordinateur, ne tombent jamais en panne. Toujours plus petit, plus fiable, plus performant, moins cher, ainsi évolue le circuit de la microélectronique. La figure F2.5 montre qu'entre 1973 et 2005, le coût de un million de transistors fabriqués sur circuit intégré est passé de 76 000 € (le prix d'un appartement) à 0,004 € (le prix d'un post-it) ! Nous sommes désormais entrés dans la nanoélectronique, et ce d'autant plus que les différentes couches dont est constitué le transistor ont des épaisseurs de quelques dizaines de nanomètres, comme on le verra par la suite concernant les composants optoélectroniques.



Figure F2.4 : Salle blanche du Laboratoire d'Architecture et d'Analyse des Systèmes

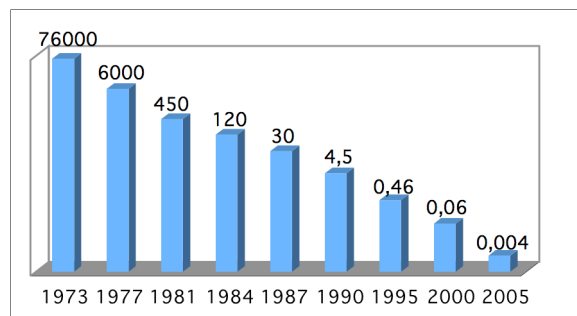


Figure F2.5 : Evolution du coût (€) de 1 million de transistors sur puce

2.2. La microélectronique à l'assaut du monde

La conséquence de cette évolution est que les circuits microélectroniques ont envahi les objets de la vie courante : ordinateurs bien évidemment, mais aussi appareils photo, caméras, téléphones, GPS, électroménagers, automobiles, pace-maker, pompe à insuline, biopuce à ADN, sécurité routière (airbag), navigation aérienne (guidage assisté), imagerie, traitement de l'information, etc., sans oublier les radars de contrôle de vitesse...

Cet essor technologique a en retour induit des retombées considérables dans tous les domaines : développement de la modélisation, des algorithmes de calcul, de démonstrations par ordinateur, avancées dans les processus de contrôle et d'automatisation, progrès immenses de l'imagerie médicale donc des connaissances médicales, pratiquement tous les domaines de la science ont été irrigués par les progrès de la microélectronique, tant il est vrai qu'une discipline ne peut jamais avancer seule bien longtemps, que la théorie et l'application se nourrissent l'une l'autre.

Je ne citerai à titre d'exemple que l'optique adaptative, récemment développée en particulier à l'observatoire de Paris. Chacun sait que les étoiles scintillent : ce scintillement est dû aux inhomogénéités de l'atmosphère, aux pollutions, aux turbulences des filets d'air chaud et d'air froid où la lumière ne se propage pas de la même manière. Il en résulte que si l'on prend deux photos à deux instants différents, elles ne sont pas identiques, ces fluctuations de la propagation lumineuse détériorent les images. Ainsi il est impossible d'obtenir de bonnes images avec un télescope à basse altitude et encore moins dans une ville. On a donc tout d'abord essayé d'installer les télescopes en altitude (Pic du Midi, Mont Palomar, télescope du Chili), où l'atmosphère est moins polluée et moins turbulente, et on les a même envoyé dans l'espace où règne le vide intersidéral (télescope Hubble). En effet, en l'absence de perturbations atmosphériques, les différents rayons lumineux issus d'une étoile parcourent le même chemin, l'onde lumineuse associée est une onde plane, et un très bon miroir (dont les aspérités doivent être inférieures au micron) donne de l'étoile une très bonne image. Mais envoyer un télescope dans l'espace constitue une solution très onéreuse.

Au contraire, les différents rayons lumineux provenant de l'étoile en traversant l'atmosphère, sont très légèrement retardés et ce de façon aléatoire en fonction des inhomogénéités de l'air dont l'indice de réfraction varie d'un endroit à l'autre (c'est ce qui cause le scintillement des étoiles, car ces inhomogénéités fluctuent sans cesse). Il en résulte que l'image de l'étoile donnée par le miroir est dégradée. L'optique adaptative repose sur un principe aussi simple que sa réalisation est complexe : il suffit d'utiliser non pas un miroir régulier (plan ou concave), mais un miroir dont la surface compense exactement les fluctuations du front d'onde, c'est à dire un miroir dont la surface fluctue en permanence et s'adapte (d'où le nom d'optique adaptative) aux fluctuations du front de l'onde lumineuse qu'il reçoit de l'étoile : ainsi les ondes incidentes fluctuent, mais les ondes réfléchies par le miroir sont planes (figure F2.6).

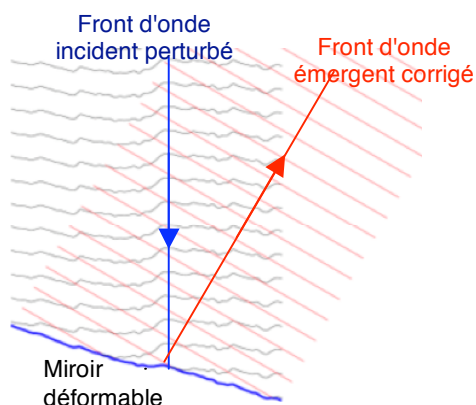


Figure F2.6 : Principe de l'optique adaptative : un miroir déformable compense les perturbations de l'onde incidente pour restituer une onde émergente plane.

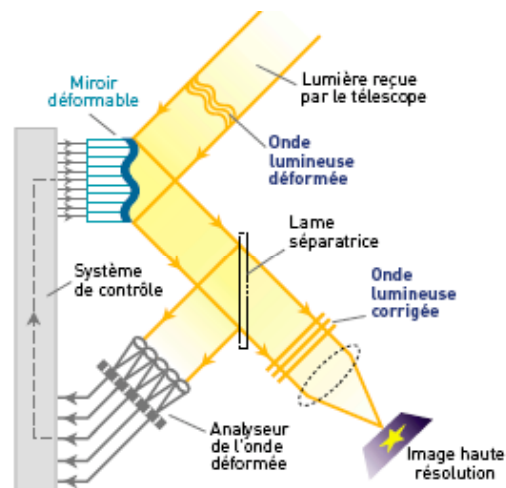


Figure F2.7 : Principe de la mise en œuvre de l'optique adaptative.

Source : "A la découverte du nanomonde", Op. cité.

En pratique (figure F2.7), le miroir du télescope est constitué d'une membrane réfléchissante flexible actionnée par un ensemble de 32 pistons dont on ajuste la hauteur au moyen de commandes électriques. Ce miroir reçoit l'onde incidente elle-même déformée, laquelle est réfléchie vers une lentille qui donne de l'étoile une image dégradée. Une partie de cette onde réfléchie est renvoyée via un miroir semi transparent (lame séparatrice) vers un analyseur d'onde, qui analyse l'onde en 32 points correspondant à la position des 32 pistons, et qui modifie en conséquence la position de chacun, jusqu'à ce que l'analyseur détecte une onde plane, alors l'image obtenue à travers la lentille est bonne.

La figure F2.8 illustre de façon saisissante l'avantage que procure ce système. Il s'agit de l'observation du centre galactique dans le proche infrarouge à une longueur d'onde de $2,2 \mu\text{m}$. A gauche, on voit l'image obtenue sans optique adaptative dans un très bon site (au Mauna Kea à Hawaï), à droite, la même image obtenue avec le système d'optique adaptative, où l'on peut même distinguer les anneaux de diffraction, ce qui montre que la résolution ultime du télescope est atteinte.

Pour que le système fonctionne, il faut que le cycle "détection d'un défaut – analyse – calcul – correction" soit beaucoup plus court que le temps moyen d'évolution d'une turbulence. Ceci n'est possible qu'avec des calculateurs ultra rapides. De tels calculateurs ne sont pas programmés mais construits spécialement pour une tâche bien définie.

Un tel programme n'a été réalisable que grâce à un concours de circonstances tout à fait remarquable :

- les ordinateurs ont acquis des puissances de calcul suffisantes : en effet tout le calcul doit être effectué au moins vingt fois par seconde. Or les 32 points d'analyse de la surface d'onde et donc les 32 points d'épreuve sur le miroir ne sont pas indépendants les uns des autres. Il faut faire intervenir l'interaction qu'ils ont les uns par rapport aux autres, ce qui augmente énormément le nombre d'éléments à calculer. Dans la pratique, c'est $32 \times 32 = 1024$ calculs qu'il faut effectuer en $1/20^{\text{ème}}$ de seconde soit 20480 calculs par seconde. Et chaque calcul porte sur 6 grandeurs différentes. C'est donc au total plus de 120 000 calculs par seconde qu'il faut effectuer.

- les techniques d'asservissement sont arrivées à un niveau de fiabilité, de rapidité, de robustesse et d'efficacité suffisant pour permettre de telles réalisations sans être à la limite des possibilités.

- enfin et surtout, les astronomes ont réussi à se faire une idée précise de ce qu'est la turbulence atmosphérique. Très schématiquement, on considère, qu'au voisinage du sol, sur une épaisseur de quelques

dizaines de mètres, l'atmosphère est constituée de globules de gaz d'environ 25 à 30 centimètres de diamètre relativement homogènes et dont la durée de vie est de l'ordre de quelques dixièmes de seconde.

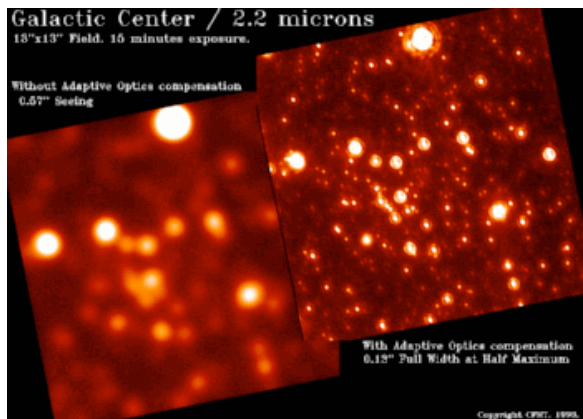


Figure F2.8 : Image du centre galactique sans (à gauche) et avec (à droite) correction par optique adaptative.
Source : L. Vapillon, J.E. Arlot, observatoire de Paris

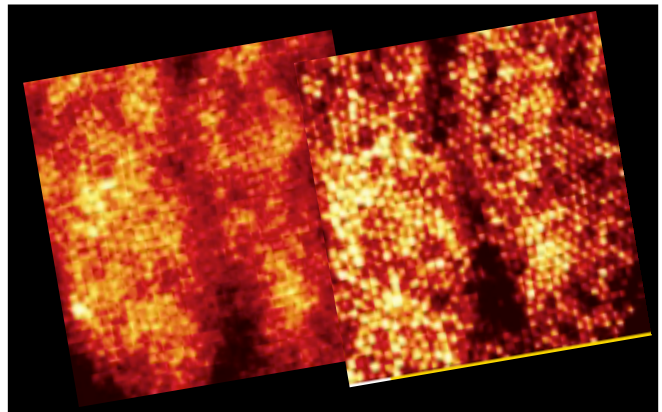


Figure F2.9 : Image de la rétine, sans (à gauche) et avec (à droite) correction par optique adaptative.
Source : A. Roorda et O. Williams, Univ. of Rochester

Cet exemple est donc une bonne illustration de l'interaction qu'ont entre elles les différentes disciplines, et du fait qu'une discipline ne peut en aucun cas progresser seule de façon spectaculaire. Une autre illustration en est donnée par les retombées dans le domaine de la santé de cette étude, initialement engagée pour les besoins de l'astronomie : c'est ainsi que l'on envisage aujourd'hui d'appliquer les techniques de l'optique adaptative aux examens ophtalmologiques. En effet, l'image de la rétine, obtenue avec les instruments traditionnels, est floue car l'humeur aqueuse de l'œil humain n'est pas homogène, tout comme l'atmosphère : l'optique adaptative permet d'obtenir des images des cellules de la rétine avec une bien meilleure netteté et un bien meilleur pouvoir de résolution, ce qui devrait permettre de détecter plus tôt et donc de mieux soigner certaines maladies rétinienues. La figure F2.9 montre deux images de la même rétine : l'image de gauche est celle obtenue directement, celle de droite a été corrigée par optique adaptative : l'amélioration de la résolution est frappante, remarquable aussi est l'analogie entre les figures F2.8 et F2.9.

Nous voici donc parvenus à l'échelle du nanomètre par la voie descendante ("top down") empruntée par l'électronique, qui a consisté à réduire progressivement les dimensions au fil des ans. Arrivés à ce stade, les lois de la physique changent, et nous devons maintenant aborder le domaine de la mécanique quantique.

3. Généralités sur la mécanique quantique

Il se trouve que les objets nanométriques échappent aux lois physiques qui gouvernent le monde à l'échelle humaine, ils obéissent aux lois de la mécanique quantique, qui décrivent des comportements tout à fait surprenants.

Par exemple, à notre échelle, nous avons conscience de pouvoir déterminer avec une précision aussi bonne que nous le voulons à la fois la position et la vitesse d'un objet : c'est ce qui fait que nous pouvons réguler un trafic ferroviaire, ou prévoir et suivre la trajectoire d'une fusée, d'un satellite, d'une planète. D'après les lois de la mécanique quantique au contraire (principe d'incertitude de Heisenberg), nous ne pouvons pas mesurer simultanément avec exactitude à la fois la vitesse et la position d'un objet de dimensions nanométriques : plus la précision sur la mesure de la vitesse est bonne, plus grande est l'erreur sur la mesure de la position, et vice-versa.

Cette approche permet de généraliser la définition de nanotechnologies, et je dirai que l'on peut qualifier de nanotechnologie toute technique qui crée des objets, ou qui étudie ou utilise des propriétés d'objets obéissant aux lois de la mécanique quantique.

Je ne vais pas ici énumérer les principes sur lesquels se fonde la mécanique quantique, ni en démontrer les conséquences. Je m'appliquerai dans les paragraphes suivants à illustrer quelques unes de ces lois en donnant des exemples concrets de réalisations qui les mettent en œuvre. On verra ainsi qu'il est

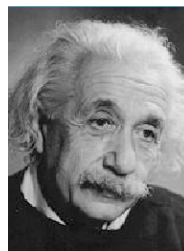
possible d'observer des nano objets, de réaliser des nano objets par diminution des dimensions (top-down) ou par agrégation d'atomes (bottom-up), d'utiliser des phénomènes quantiques dans le domaine pratique.

L'aventure quantique est jalonnée par les noms de quelques physiciens célèbres parmi lesquels : Max Planck, Albert Einstein, Niels Bohr, Louis de Broglie, Werner Heisenberg, Erwin Schrödinger, Paul-Adrien-Maurice Dirac.



Max PLANCK (1858-1947)

Pour expliquer l'émission du corps noir comportant pourtant un spectre continu, il est conduit à l'hypothèse que les niveaux d'énergie des oscillations thermiques sont quantifiés (1899) : cette année marque la date de naissance de la mécanique quantique



Albert EINSTEIN (1879-1955)

Après des études modestes, il publie en 1905-1906 quatre articles exceptionnels : dans le premier il introduit la notion de photon, grain de lumière ($E = h\nu$), dans le second il traite du mouvement brownien, dans le troisième il établit les fondements de la relativité restreinte, dans le quatrième il pose la fameuse formule $E = mc^2$.



Niels BOHR (1885-1962)

Il émet l'hypothèse que le moment cinétique de l'électron en orbite autour du noyau est un multiple entier de la constante de Planck. Sa théorie permet d'expliquer les raies d'émission de l'atome d'hydrogène. L'institut créé pour lui après la première guerre mondiale devient pendant 20 années "le phare de la physique".



Louis de BROGLIE (1892-1987)

En 1923, il émet l'hypothèse que toutes les particules se comportent comme des ondes, en particulier les électrons, théorie confirmée en 1927 par les expériences de diffraction d'électrons de Davisson et Germer.



Werner HEISENBERG (1901-1976)

Il établit en 1927 les relations d'incertitude qui portent son nom : le produit des incertitudes sur la position et sur la vitesse ne peut pas être inférieur à la constante de Planck. Aucune expérience n'a réussi à contredire ce principe.



Erwin SCHRÖDINGER (1887-1961)

Il publie en 1926 quatre articles qui jettent les fondements du formalisme de la mécanique quantique, où il établit notamment l'équation qui porte son nom, définit les énergies d'un système isolé comme étant les valeurs propres de l'opérateur hamiltonien, identifie le carré du module de la fonction d'onde à la probabilité de présence.



Paul DIRAC (1902-1984)

Il établit en 1928 l'équation qui porte son nom, et qui peut être considérée comme le fondement de la mécanique quantique relativiste. Les solutions de cette équation introduisent les notions essentielles de spin et d'antiparticule.

4. La dualité onde-corpuscule : du microscope optique au microscope électronique

4.1. Dualité onde-corpuscule

Depuis l'antiquité s'est posée la question de savoir si la lumière était constituée par une "onde lumineuse" ou par des "grains de lumière".

Les ondes se propagent, elles sont caractérisées par leur longueur d'onde λ ou par leur fréquence ν , ces deux quantités étant reliées par l'intermédiaire de la vitesse de propagation c de l'onde :

$$\lambda = c / \nu \quad (1)$$

Elles ont comme propriété caractéristique de donner lieu à des phénomènes d'interférences et de diffraction. C'est en particulier le cas pour les ondes lumineuses : la superposition d'ondes lumineuses dites cohérentes donne en certains endroits des zones d'ombre alternant avec des zones claires (franges d'interférence).



Figure F4.1 : Diffraction de la lumière par les sillons d'un disque laser

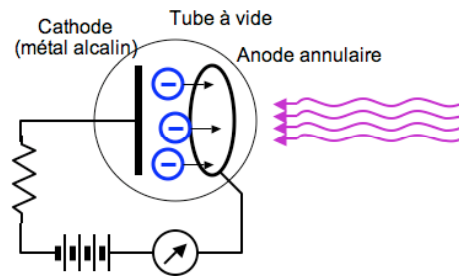


Figure F4.2 : Effet photoélectrique

On peut mettre très facilement en évidence ces phénomènes : il suffit la nuit de regarder une ampoule électrique à travers un voile aux mailles serrées : l'on voit une tâche centrale irisée et des tâches de lumière qui s'alignent le long des directions des mailles du voile. Un autre exemple spectaculaire de phénomènes d'interférences et de diffraction est constitué par la palette de couleurs magnifiques qui apparaît lorsque la lumière blanche se réfléchit sur la face striée d'un disque laser CD ou DVD (Figure 1). A la fin du 19ème siècle, la mise en évidence de ces phénomènes d'interférences et de diffraction avait apporté la preuve que la lumière n'était pas constituée de grains (c'est à dire de particules), mais au contraire se comportait comme une onde.

Or en 1887 Rudolph Hertz (célèbre par la découverte des ondes hertziennes) et son élève Philipp von Lenard découvrent l'effet photoélectrique. L'expérience est schématisée figure 2 : une tension électrique est appliquée entre deux électrodes métalliques situées dans un tube à vide (vidé de l'air qu'il contenait). De la lumière est envoyée à travers l'anode circulaire et vient frapper la cathode : l'énergie lumineuse est communiquée aux électrons situés dans la cathode, dont un certain nombre sous l'effet du choc sont arrachés au métal constituant la cathode, ils sont attirés vers l'anode et leur circulation crée un courant électrique. L'on s'attendait à ce que, lorsque l'intensité de la lumière est faible (c'est à dire l'énergie est faible), les électrons ne soient pas arrachés (pas de courant) et lorsque l'intensité lumineuse (c'est à dire l'énergie) augmente, de plus en plus d'électrons soient arrachés au métal et donc que le courant devienne de plus en plus intense. Or il n'en est rien : le courant ne dépend pas de l'intensité lumineuse, mais de la couleur de la lumière, c'est à dire de sa fréquence.

Au début du 20ème siècle, Max Planck étudie le rayonnement du corps noir et en déduit que la lumière semble arriver par paquets d'énergie, chaque paquet apportant une énergie proportionnelle à la fréquence, le facteur de proportionnalité h a depuis été appelé constante de Planck.

En 1905, Albert Einstein réalise que ces deux phénomènes, effet photoélectrique et rayonnement du corps noir, étaient deux aspects d'un même phénomène : il affirme que la lumière, considérée alors comme une onde, est aussi une particule (le photon) qui transporte une énergie :

$$E = h \nu \quad (2)$$

si la fréquence est trop faible (lumière rouge) la lumière quelle que soit son intensité ne parviendra pas à communiquer aux électrons de la cathode une énergie suffisante pour les arracher au métal, mais si la fréquence est suffisamment élevée (lumière bleue), la lumière arrachera les électrons du métal et un courant pourra circuler.

Cette théorie allant à l'encontre des idées reçues à l'époque, de nombreux physiciens tentèrent de démontrer que la théorie d'Einstein était fautive : ainsi Robert Millikan passa douze années de sa vie à réaliser des expériences pour prouver qu'Einstein avait tort, et finit par démontrer qu'il avait raison ! Millikan reçut le prix Nobel de Physique en 1923 pour avoir mesuré la charge de l'électron.

Quelques années plus tard, en 1925, Louis de Broglie émet l'hypothèse que, puisque dans certains cas une onde (lumineuse) pouvait se comporter comme des particules, réciproquement des particules devaient pouvoir se comporter comme des ondes : par exemple on doit pouvoir utiliser un faisceau d'électrons pour réaliser une expérience de diffraction : c'est effectivement ce qui put être observé 2 ans plus tard, en 1927, par Davisson et Germer. A partir d'une autre relation célèbre de Einstein :

$$E = m c^2 \quad (3)$$

De Broglie en déduit que l'on pouvait associer à l'onde une quantité de mouvement (produit de la masse par la vitesse, ici $m c$) :

$$m c = h \nu / c \quad (4)$$

Ceci établit le principe de "dualité onde-corpuscule" : à une particule caractérisée suivant la mécanique classique par son énergie E et sa quantité de mouvement $m c$, on peut associer une onde transportant une énergie E et se déplaçant à la vitesse c , de fréquence ν donc de longueur d'onde $\lambda = c / \nu$.

Finalement, la lumière est-elle une onde ou une particule (photon) ? Les deux à la fois, ou plus exactement ni l'un ni l'autre. Afin de décrire un concept, en particulier pour le transmettre à autrui, nous utilisons une représentation : les langages (français, anglais, allemand, etc.) sont des exemples de représentations, à chaque concept est associé un mot ou un ensemble de mots, qui le traduisent plus ou moins bien (d'où la difficulté de bien traduire un texte d'une langue à une autre). Les mathématiques constituent le langage qui permet le mieux de représenter les concepts de la physique : lorsque la lumière interagit avec la lumière, la représentation la mieux appropriée est en général la description au moyen d'ondes ; lorsqu'elle interagit avec la matière la description la mieux appropriée est celle qui utilise une représentation particulière (photon). Ainsi la lumière est à la fois onde et particule. En réalité, la situation est analogue à celle de l'artiste à la fois musicien et peintre, et qui traduirait ses émotions en composant une partition de musique et en peignant un tableau : les deux traduisent les émotions du peintre, mais ni le morceau de musique ni le tableau de peinture ne sont les émotions du peintre : ce sont deux représentations de ses émotions, l'une des deux étant plus appropriée que l'autre à traduire telle ou telle émotion. De même, la représentation ondulatoire et la représentation corpusculaire ne sont que deux représentations du concept de lumière, la lumière n'est ni l'un ni l'autre.

Ce qui a été dit pour l'onde lumineuse qui peut être représentée par une particule (photon) vaut aussi pour une particule tel l'électron qui dans certains cas peut se comporter comme une onde. C'est le cas plus généralement de toute onde et de toute particule.

4.2. Microscopes électroniques

Ce concept énoncé par de Broglie en 1923 allait révolutionner l'observation des petits objets. En effet, l'on sait qu'il est possible de dévier la trajectoire d'un électron en le soumettant à un champ électrique ou à un champ magnétique. Le tube cathodique de l'oscilloscope ou du téléviseur en sont des démonstrations patentes : au fond du tube cathodique du téléviseur, un canon à électrons crée un flux d'électrons que des bobines magnétiques concentrent en un point de l'écran, les déplacements de ce point sous l'effet du champ magnétique des bobines créent des images. Des lentilles optiques sont capables de faire diverger ou converger des faisceaux lumineux, et en associant de telles lentilles on peut constituer un microscope optique qui donne d'un objet une image agrandie. De la même façon, il est possible de réaliser des lentilles électrostatiques ou magnétiques qui peuvent faire diverger ou converger des faisceaux d'électrons, et en associant de telles lentilles on peut constituer un microscope électronique qui donne d'un objet une image agrandie. L'image donnée par un microscope optique est visible, celle donnée par un microscope électronique est invisible car l'onde associée aux électrons d'un tube cathodique correspond à des longueurs d'onde auxquelles la rétine n'est pas sensible. Typiquement, les longueurs d'ondes visibles sont comprises entre 0,4 et 0,8 μm soit entre 400 et 800 nm (nanomètre), alors que les longueurs d'onde associées aux faisceaux électroniques utilisées sont de l'ordre de 0,5 nm. On recueillera donc les images électroniques sur des plaques photographiques sensibles à ces longueurs d'onde ou sur des écrans fluorescents du type écran de télévision.

C'est justement cette longueur d'onde très petite qui est intéressante. Une caractéristique essentielle d'un microscope est son pouvoir séparateur : c'est la plus petite distance δ entre deux objets telles que leurs images soient séparées : si la distance entre deux points est inférieure à δ leurs images ne peuvent plus être distinguées ; ceci est dû au fait que l'image d'un point est une tâche, à cause de la diffraction ; par conséquent deux points auront comme images deux tâches ; si les deux points sont trop proches l'un de l'autre, les deux tâches se chevauchent et ne peuvent plus être distinguées. Le pouvoir séparateur δ est proportionnel à la longueur d'onde λ ; la longueur d'onde d'un faisceau d'électrons étant bien plus petite que celle d'un faisceau lumineux, le pouvoir séparateur d'un microscope électronique est bien plus petit que celui d'un microscope optique, c'est à dire qu'on peut distinguer deux points qui sont beaucoup plus proches l'un de l'autre, c'est à dire percevoir des distances bien plus petites. Typiquement le pouvoir séparateur d'un microscope optique vaut 500 nm (0,5 μm) alors que celui du microscope électronique est de l'ordre de 0,5 nm : ceci signifie que le grossissement d'un microscope électronique est 1000 fois plus grand que celui du microscope optique. Le grossissement d'un microscope électronique atteint 5 millions. Il fut inventé en 1933 par Ernst Ruska. Il atteint aujourd'hui un développement très sophistiqué, que ce soit au niveau des appareillages ou au niveau du traitement informatique des données. Les domaines explorés sont très divers : la biologie moléculaire et cellulaire, la cristallographie, la métallurgie et les sciences des matériaux.

On peut utiliser un microscope optique en transparence (c'est à dire en transmission), en étudiant la lumière qui passe au travers d'une lame mince transparente. On peut de même utiliser le microscope électronique en transmission, en recueillant le faisceau d'électrons qui traverse une lame mince, de quelques nm (lire "nanomètres") à quelques dizaines de nm. La figure F4.3.a montre la photographie prise au microscope électronique en transmission d'un empilement régulier de couches alternées de 3 nm d'épaisseur chacune de deux semiconducteurs : InAs (Arséniure d'indium, couches sombres sur la figure) et GaSb (Antimoniure de Gallium, en clair sur la figure) ; ces couches ont été déposées par épitaxie par jets moléculaires (voir ci-dessous) sur un substrat de GaSb ; on voit des lignes plus sombres à l'interface entre les couches. La figure F4.3.b constitue un agrandissement de la figure F4.3.a, réalisée au microscope électronique en transmission en mode haute résolution : l'on voit sur cette figure apparaître à la jonction entre deux couches InAs et GaSb une nouvelle couche de 0,6 nm d'épaisseur, constituée de deux couches atomiques d'un composé qui a pu être identifié comme étant de l'InSb (Antimoniure d'Indium) résultant du contact entre les atomes Sb et In respectivement de GaSb et InAs.

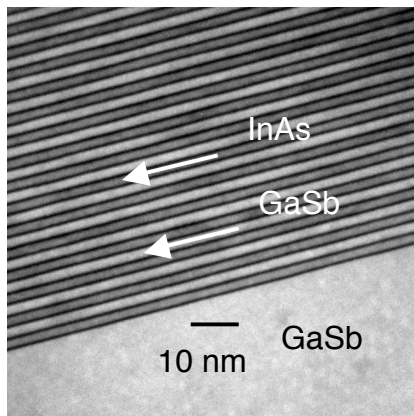


Figure F4.3.a : Photographie prise au microscope électronique en transmission d'un empilement de couches InAs-GaSb de 3 nm d'épaisseur chacune

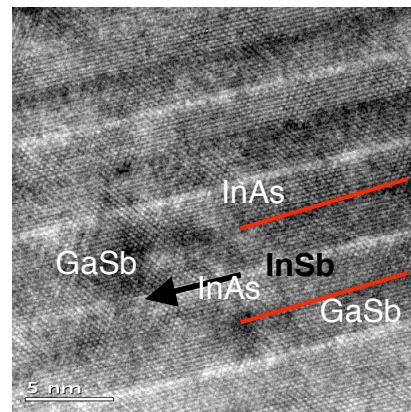


Figure F4.3.b : Microscopie électronique en transmission en mode haute résolution : identification de 2 monocouches (0.6 nm) de InSb à chaque période

Source : E. Tournié, Institut d'Electronique du Sud (UMR CNRS 5214), Université Montpellier 2.

On peut utiliser un microscope optique en réflexion : on observe la lumière réfléchie sur la surface de l'échantillon étudié. On peut aussi utiliser un microscope électronique en réflexion : dans ce cas, il s'agit de Microscopie Electronique à Balayage (MEB), en anglais SEM pour Sampling Electron Microscopy. Dans ce cas, le faisceau électronique est focalisé sur un point de la surface. Les électrons incidents (ou primaires) percutant l'échantillon lui arrachent des électrons (électrons secondaires) qui sont analysés. L'intensité de ce signal électrique dépend à la fois de la nature de l'échantillon au point d'impact qui détermine le rendement en électrons secondaires et de la topographie de l'échantillon au point considéré. Il est ainsi possible, en balayant le faisceau sur l'échantillon de façon tout à fait similaire à l'écran d'un téléviseur, d'obtenir une cartographie de la zone balayée. Lorsque les électrons incidents frappent la surface d'un échantillon : certains peuvent être réfléchis (électrons rétrodiffusés), d'autres pourront arracher des électrons aux atomes de l'échantillon (électrons secondaires), d'autres pourront perturber et exciter les électrons des couches profondes des atomes de l'échantillon : la désexcitation se matérialisera par l'émission d'électrons (électrons Auger) ou par l'émission de rayons X. L'analyse de ces différents types d'électrons et rayons X donne des informations sur la topologie de la surface et sur sa composition chimique.

5. L'effet tunnel et ses applications

5.1. Description de l'effet tunnel :

L'effet tunnel est un phénomène typiquement quantique. Sa transposition dans le monde où nous vivons serait le "passe-murailles", une personne qui pourrait traverser un mur de séparation entre deux pièces sans détériorer ni le mur ni sa personne.

En mécanique classique, si une bille qui roule sur un plan horizontal rencontre un monticule, deux cas peuvent se produire : ou bien sa vitesse (donc son énergie) est assez grande pour qu'elle "escalade" le monticule, dans ce cas elle arrive au sommet et poursuit sa route, la particule est transmise ; ou bien son

énergie est insuffisante, alors elle commence à gravir la pente, s'arrête et repart en sens inverse, la particule est réfléchiée.

En mécanique quantique, lorsqu'une onde rencontre un tel obstacle : si son énergie est assez grande, elle est en grande partie transmise mais une fraction est réfléchiée : cela signifie que si un flot de particules ayant suffisamment d'énergie frappe une marche d'escalier, la plupart continuent leur chemin mais certaines sont réfléchiées ; si l'énergie de l'onde est insuffisante, la plus grosse partie de l'onde est réfléchiée mais une partie pénètre dans la marche avant de revenir en arrière. La partie de l'onde pénétrant dans l'obstacle est l'onde évanescence. En mécanique classique, on a une onde incidente et une onde réfléchiée (figure F5.1.a), en mécanique quantique il y a une onde incidente, une onde évanescence et une onde réfléchiée (figure F5.1.b). La probabilité qu'une onde pénètre à une distance d à l'intérieur de la marche diminue lorsque d augmente.

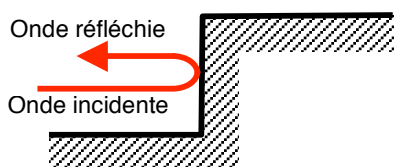


Figure F5.1.a : Mécanique classique

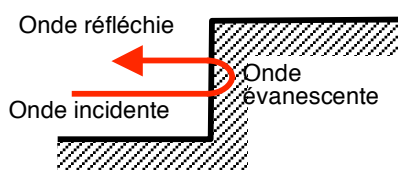


Figure F5.1.b : Mécanique quantique

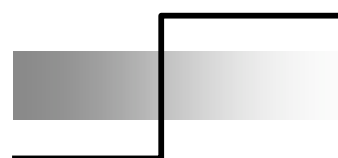


Figure F5.1.c : Représentation particulaire

Cela signifie que si l'on envoie un flot d'électrons contre la marche d'escalier, certains électrons pénètrent à l'intérieur avant de revenir en arrière, et plus on pénètre loin à l'intérieur de la marche, moins on a de chances qu'un électron puisse pénétrer jusque là. En pratique, la distance de pénétration est de l'ordre de quelques nanomètres ou dizaines de nanomètres : au delà, la probabilité de trouver une particule est si faible que l'on considère que l'on n'a aucune chance d'y rencontrer une particule. Si l'on était capable de prendre une photographie instantanée du flot d'électrons, on verrait quelque chose ressemblant à la figure F5.1.c où chaque point représente un électron, il y a autant d'électrons se dirigeant de la gauche vers la droite que d'électrons se dirigeant de la droite vers la gauche. A gauche de la paroi verticale de la figure F4.3.c se trouvent les électrons incidents et les électrons réfléchis, à droite ceux correspondant à l'onde évanescence.

Si maintenant nous remplaçons la marche d'escalier par un mur, la figure F5.1.c nous montre que, si le mur est épais (figure F5.2.a), pratiquement aucune particule n'atteint la face opposée, donc le mur n'est pas traversé. Au contraire si le mur est mince (figure F5.2.b) des particules atteignent la face opposée et donc peuvent continuer leur chemin : sans changer d'énergie, des particules peuvent traverser le mur, c'est l'effet tunnel. Le nombre de particules traversant le mur est d'autant plus grand que le mur est plus mince. En pratique, pour les énergies habituelles rencontrées en électronique, l'effet tunnel se manifeste pour des électrons pour des épaisseurs inférieures à quelques nanomètres.



Figure F5.2.a : Une barrière épaisse n'est pas franchie



Figure F5.2.b : Franchissement d'une barrière mince par effet tunnel.

5.2. Microscope à effet tunnel :

Un nouveau type de microscope, le microscope à effet tunnel (en anglais STM, Scanning Tunneling Microscope) est inventé en 1981 par des chercheurs d'IBM, Gerd Binnig et Heinrich Rohrer, qui reçurent le Prix Nobel de physique pour cette invention en 1986. Cette nouvelle technique utilisant une pointe très fine terminée par un atome permet l'observation directe et aisée d'atomes et de structures atomiques de surfaces conductrices dans une large variété d'environnements (air, eau, huile, vide).

Le principe de fonctionnement est le suivant (voir figure F5.3) : Il s'agit, pour simplifier, d'un palpeur (une pointe) qui suit la surface de l'objet. Pour cela, avec un système de positionnement de grande précision (réalisé à l'aide de capteurs piézoélectriques), on place la pointe conductrice en face de la surface à étudier et l'on mesure le courant résultant du passage d'électrons entre la pointe et la surface par effet tunnel. Ce courant varie de façon exponentielle avec la distance de la pointe à la surface de l'échantillon, avec une

distance caractéristique de quelques dixièmes de nanomètres. Puis on déplace la pointe du levier le long de l'échantillon (balayage) et on ajuste sa hauteur de manière à conserver une intensité du courant tunnel constante. On peut alors déterminer le profil de la surface avec une précision inférieure aux distances interatomiques. On reconstitue ainsi non pas une photographie, mais une image de synthèse de la surface c'est à dire des atomes qui la composent. Il faut bien entendu que la pointe soit de l'ordre d'une fraction de nanomètre. Afin d'obtenir une bonne précision, la distance entre la pointe et la surface ne doit pas fluctuer de plus d'une fraction de nanomètre (typiquement 0,1 nm), d'où la nécessité absolue que l'appareil soit protégé des vibrations. La figure F5.4 montre une vue d'un microscope à effet tunnel.

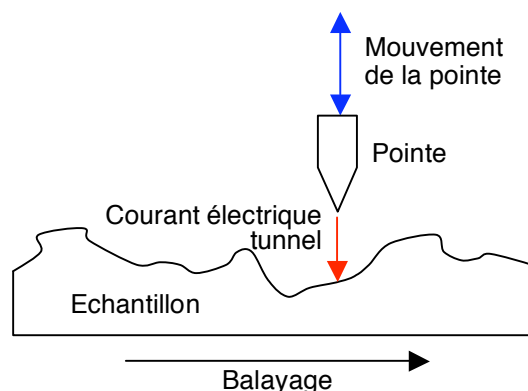


Figure F5.3 : Principe de fonctionnement du microscope à effet tunnel

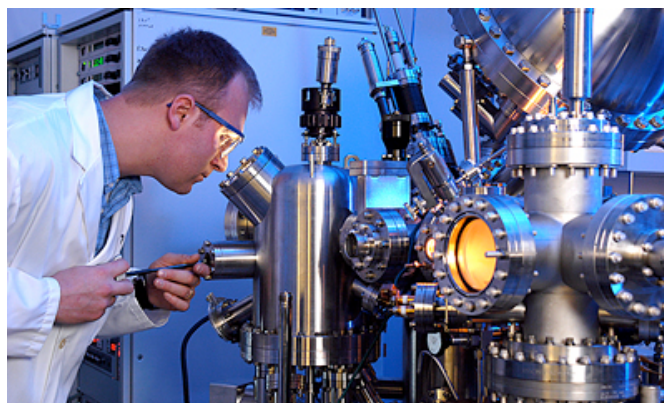


Figure F5.4 : Microscope à effet tunnel
Source : Conseil National de la Recherche du Canada

En réalité le schéma de fonctionnement même dans son principe n'est pas aussi simple que le montre la figure F5.3, il ressemble plutôt à la figure F5.5 : la pointe est constituée d'atomes (figure F5.5.a), plus ou moins régulièrement disposés, et lorsqu'on la déplace le long d'un échantillon en maintenant le courant tunnel constant, donc la hauteur de la pointe par rapport à l'échantillon constante, on obtient un tracé qui ressemble à la figure F5.5.b : ayant obtenu ce tracé, on doit ensuite traiter cette image pour reconstituer le profil de l'échantillon.

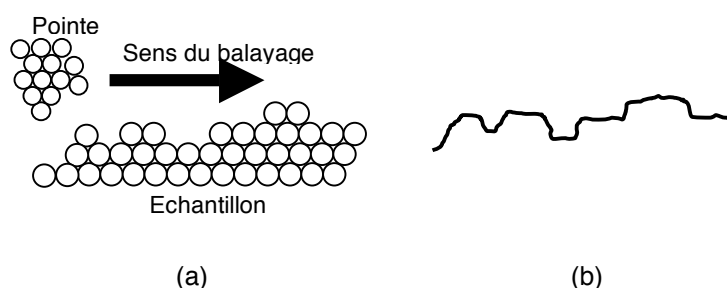


Figure F5.5 : Tracé obtenu (b) à partir du balayage de la surface de l'échantillon (a)

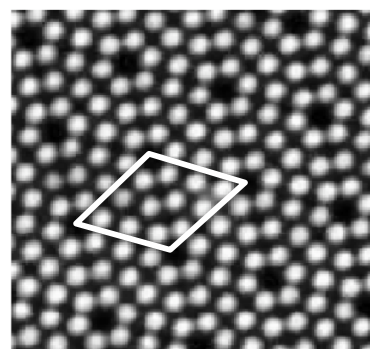


Figure F5.6 : Structure de surface du Silicium (D. Rodichev et J. Klein, 227^e Conférence de l'UTLS, Août 2000)

La figure F5.5 illustre l'importance du traitement de l'information, qui ne cesse de progresser concernant aussi la microscopie à effet tunnel. On enregistre dans l'expérience un relief en fonction des coordonnées d'un point quelconque de la surface. Ce signal électronique est traité à l'aide d'un ordinateur et transformé en nuances de gris ou de couleurs. En général, les points les plus hauts sont blancs, les plus bas noir et les points intermédiaires convertis en nuances de gris. Différentes représentations sont alors possibles, vue de dessus, perspectives, couleurs, etc. La figure F5.6 montre surface du silicium, faisant apparaître une structure (7x7), différente de la structure en volume et longtemps demeurée inconnue, où la maille que l'on doit répéter pour constituer la surface périodique est un losange (tracé en blanc sur la figure F5.6) de 2,8 nm de côté possédant aux quatre coins un trou en noir sur la figure, les boules blanches représentent des atomes de silicium.

On peut aussi effectuer un traitement en fausses couleurs. La figure F5.7 montre une image STM à courant constant de la surface de Nickel (111), au voisinage d'un défaut ; la taille de l'image est 2,3 nm x 2,3 nm ; la tension tunnel appliquée vaut 10 mV et le courant tunnel 5 nA.

Il est important de noter que les figures F5.6 et F5.7 ne constituent pas des photos d'atomes : il s'agit comme on vient de le voir d'images de synthèse, de reconstitution topographique de surface, c'est à dire dans le meilleur des cas de la surface en dessous de laquelle la pointe ne peut pas pénétrer. Chacune des boules représentées sur ces images ne représente pas la forme d'un atome, mais la "sphère d'influence" de l'atome. On sait que l'atome, quant à lui, est fait de vide.

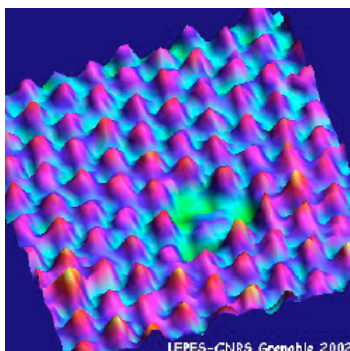


Figure F5.7 : Image au microscope à effet tunnel de la surface du Nickel (orientation 111). Source : P. Mallet et J.Y. Veuillen, LEPES-CNRS, Grenoble

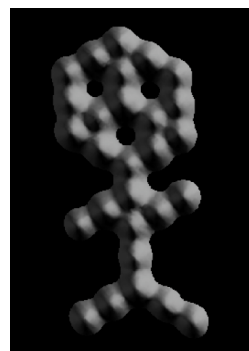


Figure F5.8 : "COman" réalisé par STM (D.M. Eigler de IBM Almaden Research Center aux Etats Unis ; <http://www.almaden.ibm.com/vis/stm/>)

Le microscope à effet tunnel permet l'étude :

- des surfaces des semi-conducteurs et des métaux, ce qui nécessite fréquemment d'observer ces surfaces dans un environnement ultravide, de l'ordre de 10^{-7} à 10^{-8} torr (ce qui revient à diluer un cube d'air de 0,5 cm de côté dans 1 m³ de vide),
- des molécules absorbées sur une surface, molécules organiques ou d'intérêt biologique.

On peut à l'aide d'un microscope à effet tunnel piloter une réaction chimique locale : il est possible d'accrocher avec la pointe un atome, de le déplacer et de le lâcher sur un site où on désire le faire réagir. La figure F5.8 montre ainsi un "COman", un bonhomme dessiné avec des molécules de monoxyde de carbone CO déposées sur un substrat à l'aide d'un microscope à effet tunnel.

Le microscope à effet tunnel est déjà entré dans le domaine des applications industrielles, il sert aujourd'hui de guide pour la fabrication des réseaux de diffraction, pour la caractérisation des têtes de lecture de films magnétiques ainsi que les films eux-mêmes, pour le contrôle de fabrication de disques compacts, etc.

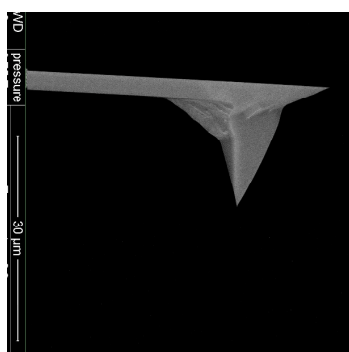


Figure F5.9 : Photographie de l'extrémité d'un levier : l'image est un carré de 50 μm de côté (Photo réalisée par P. Girard, IES, Montpellier)

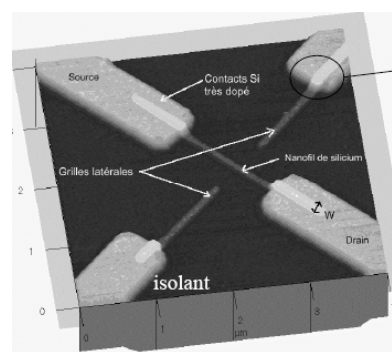


Figure 5.10 : Transistor construit par AFM. L'image a une dimension de 4 μm x 4 μm . Source : F. Marchi : "Microscopies en champ proche : techniques d'exploration et outil de manipulation et de fabrication dans le nanomonde" (LEPES-CNRS, Grenoble)

5.3. Microscope à Force Atomique (AFM) :

Le microscope à effet tunnel ne permet pas d'imager des surfaces isolantes puisqu'alors un courant électrique ne peut s'établir entre la pointe et l'isolant. En 1986, G. Binnig, C.F. Quate et C. Gerber inventèrent le microscope à force atomique. La cohésion des atomes d'un matériau est assurée par les

forces interatomiques : si l'on écarte les atomes, une force d'attraction les rapproche, si on les rapproche, une force de répulsion les écarte : l'équilibre entre atomes du matériau est réalisé lorsqu'il y a égalité entre les forces d'attraction et de répulsion. Le principe du microscope à force atomique consiste à mesurer la force d'interaction entre l'atome du bout de la pointe et les atomes de la surface de l'échantillon. La pointe est solidaire d'un bras de levier (figure F5.9) et l'échantillon est ici déplacé devant la pointe. En mesurant la déflexion du bras de levier, on obtient une mesure directe de la force pointe-substrat. La déflexion du bras de levier est en général mesurée par la déflexion d'un faisceau laser, réfléchi par un miroir monté à l'arrière du bras de levier. Des forces aussi faibles que 10^{-9} N (N = Newton) ont ainsi pu être détectées (soit le poids d'un gouttelette d'eau de trois centièmes de millimètres de rayon).

Le microscope AFM permet aussi de réaliser des motifs par assemblages d'atomes. La figure F5.10 montre un transistor assemblé sur un substrat, l'image a $4\text{ }\mu\text{m}$ de côté. La source est en haut à gauche de l'image, le drain en bas à droite, tous deux sont reliés par un fil de silicium déposé sur le substrat, qui constitue le canal, et dont la largeur est de quelques nanomètres. Deux grilles en silicium également sont déposées près du canal afin de contrôler le courant. Le courant de saturation de ce transistor est de l'ordre de $2\text{ }\mu\text{A}$ pour une tension drain de 5 V pour une tension de la grille latérale de 5 V.

Par rapport au paragraphe 2 ci-dessus, on voit que l'on est passé des nanotechnologies par voie descendante (top-down) au stade industriel, à une nanotechnologie ascendante (bottom-up) artisanale.

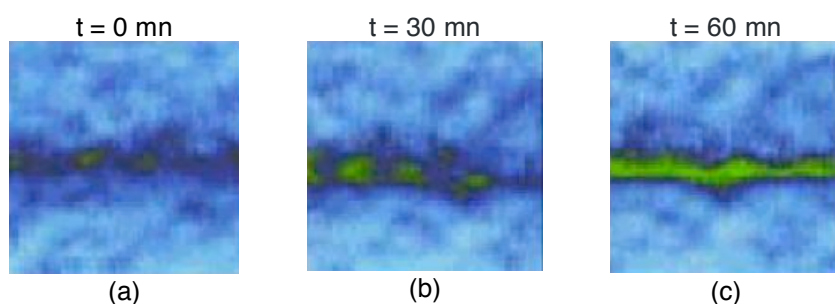


Figure F5.11 : Image en microscopie à force atomique de la constitution d'une fracture dans un verre amorphe par coalescence de nanocavités. Chaque image mesure 75 nm de côté. (Source : M. Ciccotti, M. George, A. Grimsaldi, LCVN, Univ. Montpellier 2, UMR CNRS 5587)

La figure F5.11 montre la constitution d'une fracture dans du verre amorphe, donc isolant électrique, visualisée par microscopie à force atomique : ces images prises à 30 minutes d'intervalle montrent un front de fissure se propageant de gauche à droite à une vitesse moyenne de l'ordre de 10^{-11} m/s : (a) : mise en évidence de cavités nanométriques devant la tête de fissure ; (b) : croissance des cavités ; (c) : La fracture avance par coalescence de ces nano-cavité.

Un champ d'application privilégié de la microscopie à force atomique concerne la biologie, car les objets biologiques, sont en général isolants. L'une des applications les plus fascinantes dans ce domaine est l'étude in vitro de l'ADN et des interactions ADN-protéines. Faire une image topographique par balayage d'une surface nécessitait environ une minute pour tous les microscopes existant à la fin des années 1990. Aujourd'hui, les nouveaux microscopes tunnels ou à force atomique permettent de faire une image en moins d'une seconde. Si un phénomène relativement lent se produit sur la surface examinée, il pourra donc être pratiquement filmé en prenant des images successives toutes les secondes par exemple. On peut ainsi filmer les déplacements d'objets microscopiques sur les surfaces. A titre d'exemple, on a pu filmer l'entrée d'un virus dans une cellule.

Depuis l'AFM, d'autres types de microscopes dits "en champ proche" ont vu le jour : le microscope à force magnétique, qui permet de déterminer la structure magnétique d'un échantillon ; le microscope à conductance ionique, qui n'a pas la résolution atomique mais qui permet de visualiser en milieu liquide la surface d'une membrane biologique ; le microscope optique à onde évanescente.

6. Conclusion

Dans cette conférence j'ai passé sous silence un nombre considérable de nano objets et de nano applications, pourtant importantes et/ou prometteuses : nanotubes de carbone, membranes nanoporeuses, nanoparticules utilisées dans les cosmétiques, biopuces, nanocristaux fluorescents, poussières électroniques, implants biocompatibles, électronique moléculaire, vecteurs ciblés de médicaments, etc. Cependant les quelques exemples qui ont été donnés présentent un certain nombre de points communs très généraux et permettent de dégager quelques grands axes de réflexion :

– pratiquement aucun des objets décrits n'existaient il y a à peine 50 ans : il est maintenant banal de dire que ces décennies ont connu un essor technologique sans précédent qui a permis de faire des progrès fulgurants, la microélectronique en est un exemple caractéristique,

– si l'homme est très loin d'être en mesure d'imiter la nature, il est cependant maintenant capable de créer des structures totalement artificielles de dimensions nanométriques possédant des fonctionnalités très intéressantes : hétérostructures, nanotubes de carbone, etc.,

– dans ce domaine comme dans beaucoup d'autres, il y a une grande interdépendance entre la recherche fondamentale et les applications. L'on constate en fait d'une part que les applications ne peuvent pas voir le jour si les études fondamentales ne sont pas suffisamment avancées (sans la mécanique quantique, les nanotechnologies n'existeraient pas), d'autre part que la théorie ne peut pas se développer sans les avancées technologiques : lors de la Conférence des Semiconducteurs de Boston en 1971, Pierre Aigrin pionnier de la physique des solides et ancien ministre déclarait qu'il n'y avait plus grand chose à découvrir en physique des solides : l'avènement de la microélectronique et le foisonnement de recherches fondamentales suscitées par les propriétés nouvelles liées à la diminution des dimensions lui ont donné un cinglant démenti ; l'essor de l'informatique en particulier fondamentale, le renouveau des mathématiques (par exemple dans la théorie et le traitement du signal, la cryptographie, etc.) doivent énormément à l'augmentation de puissance des ordinateurs ; l'on pourrait multiplier les exemples à l'infini. La séparation idéologique, qui date de plus d'un siècle, entre science fondamentale et science appliquée apparaît de tout évidence aujourd'hui comme une querelle dépassée. Fort heureusement, le mur de Berlin que les puristes avaient construit autour de la science fondamentale est en train de s'écrouler.

– une seconde évolution s'est dessinée depuis une vingtaine d'années, et se confirme de jour en jour : c'est le développement et la nécessité de l'interdisciplinarité : aujourd'hui pratiquement aucune science ne peut avancer seule, tous les domaines sont interdépendants. Il est bien évident que les progrès récents considérables de la médecine n'auraient pas vu le jour sans les progrès de la génétique, de l'analyse chimique, de l'imagerie médicale, cette dernière technique à elle seule nécessitant des appareillages sophistiqués, la compréhension de phénomènes physiques, des méthodes élaborées de traitement d'image et de reconstruction tridimensionnelle, des ordinateurs puissants et donc de l'électronique de commande et de la microélectronique, etc. Même les mathématiques doivent sortir de leur tour d'ivoire, sollicitées et interpellées par les autres disciplines, et dont le questionnement génère de nouvelles recherches en mathématiques sur des problèmes de modélisation, de cryptographie, etc. Même les sciences humaines et sociales sont interpellées, entre autres par des problèmes de linguistique, de sémantique, de cartographie pour les géographes, d'accès à l'information et à des bibliothèques numérisées. Cette interdisciplinarité n'exige pas que les électroniciens fassent de la biologie et que les biologistes fassent de l'électronique, elle exige au contraire des disciplines fortes et des équipes interdisciplinaires travaillant ensemble en interaction forte pour qu'en particulier chacun apprenne le langage de l'autre.

– après l'avènement des nanosciences et nanotechnologies, leur développement apparaît maintenant inéluctable, car il s'inscrit dans le cadre d'enjeux sociétaux, économiques, et stratégiques :

- sociétaux par exemple par l'impact sur le public et sur la mentalité des citoyens de l'informatisation de la société, et du développement des communications : nous ne travaillons plus et nous ne vivons plus de la même façon depuis que l'ordinateur est devenu un outil banalisé de tout un chacun, depuis que l'on communique depuis chez soi avec la terre entière en utilisant le web, depuis qu'on peut appeler n'importe qui et n'importe où avec un téléphone portable ;
- économiques : nous avons vu l'énorme marché que représentent les diodes électroluminescentes, les têtes de lecture de CD et DVD, les appareils d'imagerie numérique médicale, les télécommunications, les ordinateurs. On sait que les technologies de l'information et de la communication ont contribué pour 1/3 de la croissance des Etats Unis et pour 1/4 de la croissance de la France entre 1995 et 2000 ;
- stratégiques : le monopole ou au moins la suprématie dans les domaines des télécommunications, du cryptage, des grands ordinateurs, offre à qui les possède une position stratégique dans le monde.

– la conséquence en est qu'un effort de recherche important et soutenu doit être consenti dans les domaines couverts par les nanotechnologies. Le sujet est si vaste et si important qu'un pays seul tel que la France ne peut en aucun cas affronter tous les défis, les enjeux sont au moins à l'échelle d'un continent, et la puissance publique tout autant que le secteur privé doivent, chacun dans le cadre de ses compétences propres, apporter leurs contributions.

– enfin l'on ne saurait conclure sans mentionner le fait que, comme tous les nouveaux champs d'investigation, les nanotechnologies sont à la fois porteuses de grands espoirs et génératrices de risques et de dangers nouveaux. Ce chapitre à lui seul pourrait faire l'objet d'une conférence. Pour rester schématique, nous dirons que les dangers sont de deux sortes :

- les risques technologiques : ainsi par exemple les nanotubes de carbone, présentent des risques de pollution. Compte tenu de leur petite taille, les nanotubes peuvent facilement être absorbés par l'organisme, et compte tenu de leur caractère de cycle benzénique polymérisé, la question de l'intercalement entre les cycles d'ADN et des risques élevés de cancer en résultant sont objet à interrogations. Leur impact sur l'environnement ainsi que sur la santé fait actuellement l'objet d'études. Un article du journal *Langmuir* de l'American Chemical Society a récemment étudié le caractère "tueur de cellules" des nanotubes par contact direct en déchirant les membranes cellulaires. D'une manière générale, la matière fractionnée ou organisée à l'échelle du nanomètre d'une part présente une réactivité beaucoup plus grande que le même matériau dans les conditions habituelles, d'autre part possède de par sa petite taille un pouvoir de pénétration et d'infiltration dans l'organisme auxquels il faut prêter une grande attention, si l'on ne veut pas que les nanoparticules deviennent l'amiante du 21ème siècle.
- les risques sociétaux : la miniaturisation des composants fait que les espaces publics ou privés peuvent être truffés de systèmes de surveillance pouvant risquer de porter atteinte à la vie privée. Les moyens importants de surveillance des communications, l'organisation des systèmes d'information et la constitution de bases de données gigantesques font que chaque citoyen peut être fiché. Certes les états policiers et répressifs n'ont pas attendu l'avènement de l'informatique et des nanotechnologies pour sévir, mais ce qu'il a été possible de réaliser avec des moyens artisanaux risque de prendre une ampleur bien plus considérable avec les moyens dont nous disposons actuellement, ce qui nécessite par conséquent une vigilance encore accrue pour éviter de telles déviations.

Quoi qu'il en soit, après avoir constaté que nous baignons dans le nanomonde des molécules naturelles comme monsieur Jourdain faisait de la prose, les exemples développés dans la présente conférence vous auront convaincu je pense, que nous avons commencé depuis deux décennies déjà à baigner dans le nanomonde des objets artificiels que l'être humain est désormais capable de fabriquer. Ce phénomène ne pourra que s'amplifier dans les 20 ou 30 prochaines années.