

LA MODÉLISATION DES SURFACES, INTERFACES ET NANOSTRUCTURES

La compréhension et la maîtrise de processus comme la nanostructuration spontanée des surfaces ou la conduction entre une pointe microscopique et une surface sont essentielles dans l'industrie des nouvelles technologies, où la miniaturisation est un enjeu majeur. Dans ce domaine, la modélisation doit prendre en compte le caractère atomique de la matière avec une très grande précision.

Transfert d'échantillon dans un microscope à effet tunnel permettant l'étude de nanostructures en matériaux semiconducteurs au centre CEA de Grenoble.



A. Gonin/CEA

La physique des surfaces et interfaces est un domaine très vaste de la physique de la matière condensée. Son champ d'investigation regroupe l'étude de phénomènes aussi quotidiens que l'étalement d'une goutte sur une surface solide ou la corrosion d'un matériau, mais aussi de processus moins familiers tels que la nanostructuration "spontanée" d'une surface cristalline ou la conduction entre une pointe métallique microscopique et une surface.

La compréhension précise des phénomènes de surface nécessite une description à l'échelle atomique qui fait toute la difficulté de leur **modélisation**. Les modèles théoriques doivent prendre en compte le caractère atomique de la matière avec une précision suffisante. Trois exemples d'études effectuées au CEA au sein du groupe MSIN (Modélisation des surfaces, interfaces et nanostructures) du Service de physique et chimie des surfaces et interfaces illustrent la méthodologie employée. Le premier traite du problème de la forme d'équilibre d'une surface métallique,

le deuxième de l'imagerie des surfaces à l'échelle atomique et le troisième de la conduction électrique le long des interfaces conducteur-isolant-semiconducteur⁽¹⁾ utilisées dans l'industrie microélectronique.

Jouer de l'instabilité de certaines surfaces vicinales

Dans le premier exemple, il s'agit d'étudier le profil d'équilibre d'une surface. En effet, selon leur morphologie, les surfaces peuvent constituer un support approprié pour l'élaboration de divers types de **nanostructures**, aux multiples applications potentielles. Le point de départ est une surface bien particulière dite "vicinale", obtenue en coupant un cristal par un plan cristallographique faiblement désorienté par rapport à un plan de grande densité atomique. Une telle surface se présente comme une succession périodique de terrasses séparées par des marches de hauteur monoatomique (figure 1a) qui présentent des caractéristiques intéressantes. D'une

part, ces terrasses sont de dimension nanométrique⁽²⁾ et, d'autre part, il est possible de tirer parti d'une instabilité naturelle de certaines de ces surfaces qui ont tendance à adopter spontanément un profil en "toit d'usine" (figure 1b) tout en conservant leur orientation moyenne. Ce phénomène est connu sous le nom de facettage. Les atomes déposés sur ces surfaces vont préférentiellement occuper les sites se trouvant dans les angles rentrants, où leur nombre de liaisons sera maximum. Il est ainsi possible de construire un réseau

(1) Un semiconducteur est un matériau dans lequel la bande d'états électroniques occupés (bande de valence) est séparée de la bande des états inoccupés (bande de conduction) par une bande d'énergie interdite relativement étroite. Un tel matériau est un isolant électrique au zéro absolu, mais devient modérément conducteur lorsque sa température est suffisamment élevée pour exciter des électrons de la bande de valence vers la bande de conduction.

(2) 1 nanomètre = 10^{-9} m.

Figure 1a. Surface vicinale à marches monoatomiques.

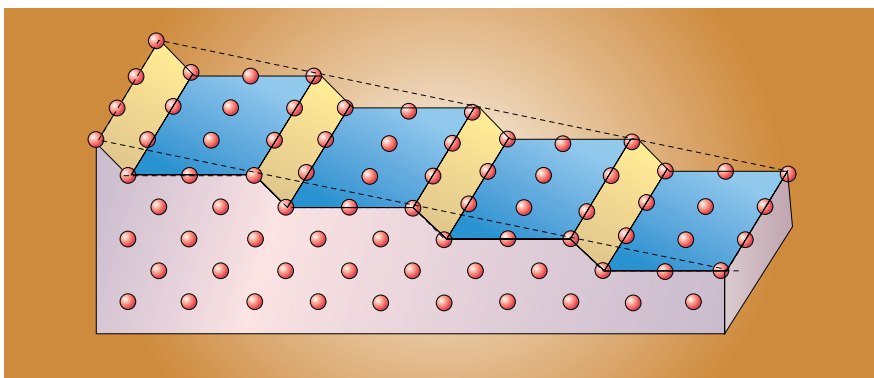
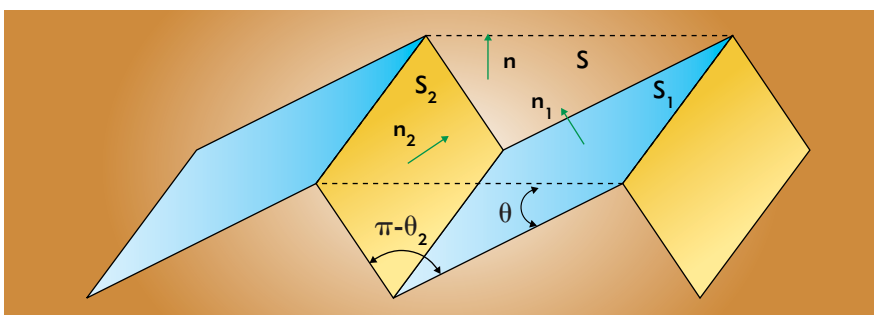


Figure 1b. Surface facettée en "toit d'usine".



Le facettage en formules

1

Une aire S d'une surface vicinale se transforme en facettes de normale n_1 (aire S_1) et de normale n_2 (aire S_2) si la transformation abaisse l'énergie du système. En désignant par γ , γ_1 et γ_2 les énergies de surface par surface unité des surfaces correspondantes, cela s'exprime sous la forme :

$$\gamma_1 S_1 + \gamma_2 S_2 < \gamma S$$

avec la conservation de l'orientation moyenne pour contrainte, soit :

$$S = S_1 \cos \theta + S_2 \cos (\theta_2 - \theta), S_1 \sin \theta = S_2 \sin (\theta_2 - \theta)$$

où θ est l'angle entre S et S_1 , et θ_2 l'angle entre les deux facettes S_2 et S_1 . Ces conditions peuvent être réunies en une seule :

$$\gamma / \cos \theta > (1 - \tan \theta / \tan \theta_2) \gamma_1 + (\tan \theta / \tan \theta_2) \gamma_2 / \cos \theta_2$$

Cette inégalité a une signification géométrique simple : la surface d'orientation θ est instable si son point représentatif dans un diagramme $\gamma / \cos \theta = f(\tan \theta)$ est au dessus de la droite joignant les points de coordonnées $(0, \gamma_1)$ et $(\tan \theta_2, \gamma_2 / \cos \theta_2)$ (figure) et stable dans le cas contraire.

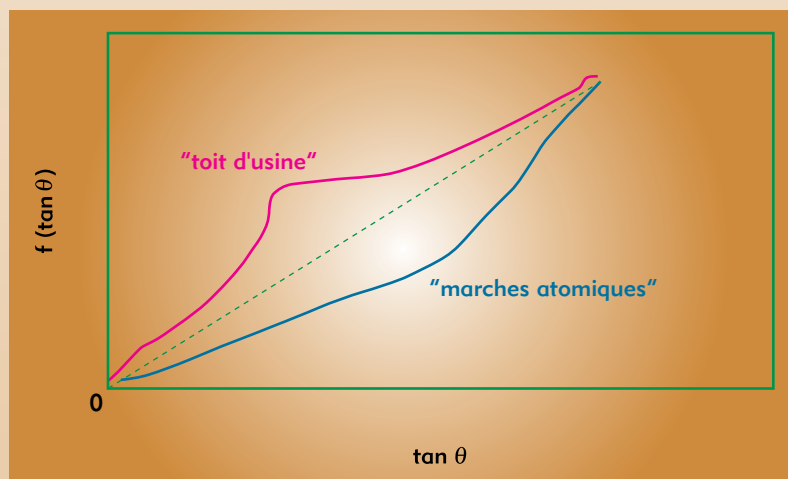


Figure. Représentation schématique de la fonction $\gamma / \cos \theta = f(\tan \theta)$ et de la condition de stabilité (ou instabilité) de la surface vicinale.

périodique de **nanofils** dont les propriétés magnétiques et de transport électronique sont d'un grand intérêt technologique.

La condition de facettage (encadré 1) implique de pouvoir calculer l'énergie de surface, c'est-à-dire l'énergie nécessaire pour créer une surface d'aire unité pour une orientation quelconque. Cette énergie est liée à la rupture de liaisons chimiques lors de la création de la surface. Des diverses approches pour la calculer, la plus simple est basée sur l'utilisation de **potentiels empiriques** qui décrivent l'interaction mutuelle des atomes, un peu comme les planètes interagissent en fonction de la distance qui les sépare. La deuxième, beaucoup plus complexe, part de la description **quantique** des phénomènes mis en jeu. Il a été prouvé que l'utilisation de potentiels empiriques conduisait à un comportement très réducteur (illustration page suivante), car beaucoup trop schématique, et que seul un calcul quantique complet (encadré 2) permet de déterminer si une surface vicinale est stable ou instable par rapport au facettage. Cet important résultat montre, une fois de plus, que les propriétés des surfaces sont gouvernées par la physique et la chimie des liaisons atomiques, dont la description nécessite des modèles élaborés.

Voir et comprendre les surfaces à l'échelle atomique

Le deuxième exemple porte sur la vision et, mieux encore, la compréhension des surfaces à l'échelle atomique, rendues possibles par le microscope à effet tunnel (ou STM). Depuis son apparition aux début des années 1980, le STM a permis aux scientifiques en physique et chimie des surfaces de réaliser un vieux rêve : comprendre et "manipuler" la



Les électrons dans les solides

2

Les électrons sont à l'origine des propriétés de cohésion et de conduction électrique des matériaux. Ils n'obéissent plus aux lois de la mécanique classique mais à celles de la mécanique **quantique**. Leur comportement n'est pas décrit par une trajectoire mais par une fonction d'onde complexe $\Psi(r)$ régie par une **équation aux dérivées partielles** $H \Psi(r) = E \Psi(r)$, appelée **équation de Schrödinger**, qui donne également l'énergie E de la particule. L'opérateur H , appelé hamiltonien, est la somme des opérateurs énergie cinétique et énergie potentielle des électrons. Le carré de la fonction d'onde $|\Psi(r)|^2$ donne la probabilité de présence de l'électron en un point r de l'espace, ce qui implique que $\Psi(r)$ satisfasse à certaines conditions pour être physiquement acceptable. Pour les électrons liés à un atome, ces conditions ne peuvent être réalisées que pour des valeurs discrètes des niveaux d'énergie qui sont dits "quantifiés" (d'où le nom de *mécanique quantique*). Ces électrons occupent les niveaux par ordre d'énergie croissante à raison de 2 au maximum par niveau. Si l'on rapproche les atomes pour former un solide, les niveaux "de cœur" correspondant aux électrons les plus proches du noyau sont peu perturbés, par contre les électrons de **valence** qui forment les couches électroniques externes, d'énergie plus élevée, se délocalisent ce qui assure la cohésion du solide. Leurs niveaux forment des **bandes d'énergie** (figure), remplies jusqu'au dernier

niveau occupé appelé *niveau de Fermi*. Une composante essentielle de l'énergie totale d'un système est obtenue en faisant la somme des énergies de tous ses électrons. La modification des niveaux d'énergie et des fonctions d'onde au voisinage d'une surface est à l'origine de l'énergie de surface. De plus, en présence d'une surface, si les électrons restent essentiellement confinés dans le solide par une barrière de potentiel de hauteur finie, ils ont cependant une probabilité non nulle de s'étendre dans le vide. Ce sont ces électrons qui sortent de la surface et dont l'énergie est proche du niveau de Fermi qui sont détectés par l'établissement d'un courant électrique (dit "courant tunnel") en approchant une pointe métallique à une distance de l'ordre du nanomètre⁽²⁾.

La détermination théorique de toute propriété physique du solide passe donc par la résolution de l'équation de Schrödinger qui, malgré son apparence simple, ne peut être effectuée que par de longs calculs sur ordinateur, à partir d'hypothèses simplificatrices mais raisonnables. Dans certains cas, il existe des potentiels classiques d'interaction entre atomes, plus ou moins empiriques, qui permettent de calculer l'énergie du système. Ces potentiels contiennent de manière implicite et approchée l'effet des électrons mais sont trop rudimentaires pour rendre pleinement compte de leur caractère quantique.

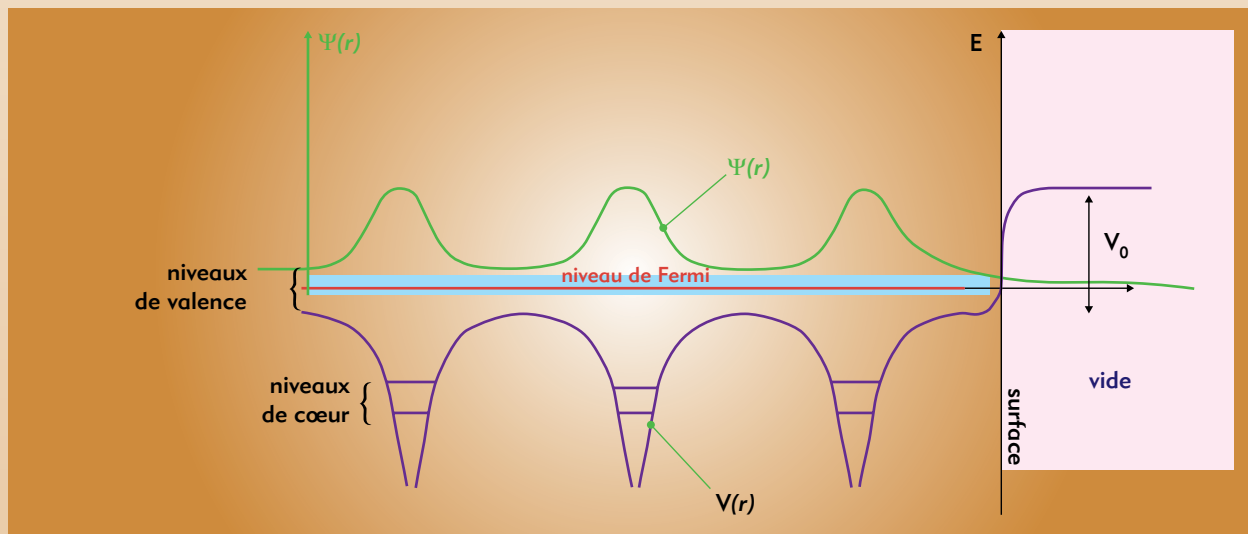
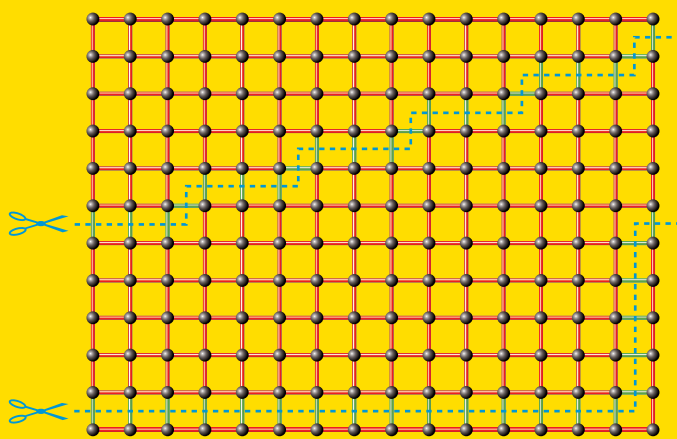


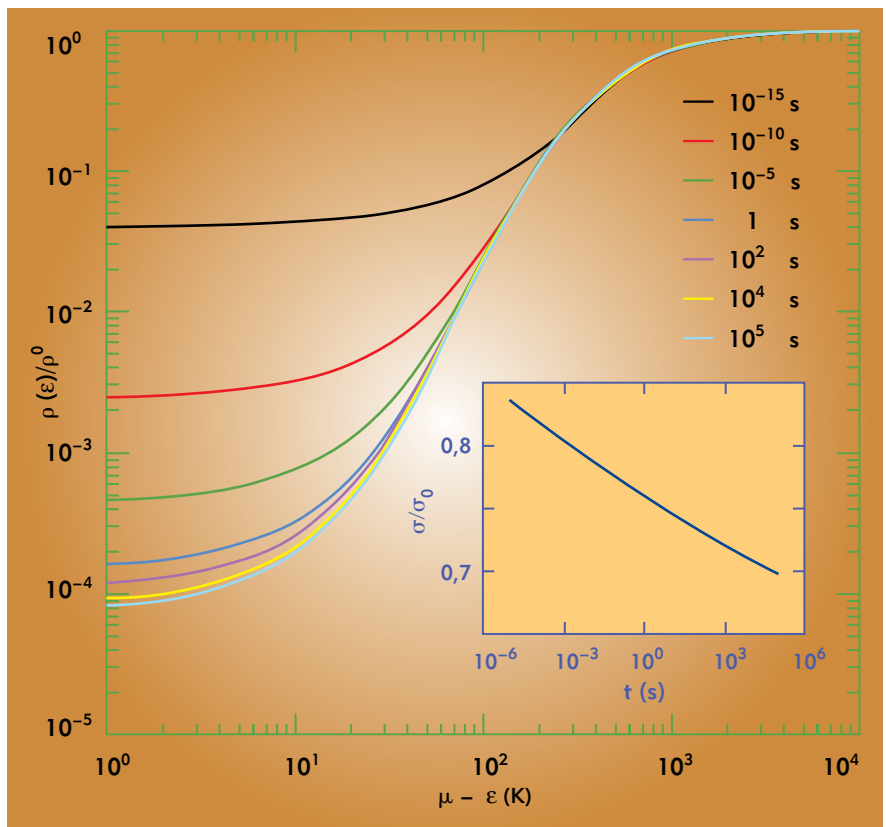
Figure. Représentation schématique du potentiel $V(r)$ vu par un électron et de sa fonction d'onde $\Psi(r)$ au voisinage d'une surface caractérisée par sa barrière de potentiel V_0 . Chaque puits de potentiel correspond à un atome.



● ● ● ● ●

Modèle énergétique le plus simple pour décrire un cristal : l'énergie totale s'écrit comme une somme d'interactions entre paires d'atomes premiers voisins. Dans ce modèle, l'énergie de surface est déterminée par le nombre de liaisons coupées (en vert). Il apparaît immédiatement qu'avec ce type de potentiel empirique, la surface vicinale et la surface facettée ont exactement la même énergie (même nombre de liaisons coupées), ce qui ne permet donc pas de résoudre le problème de la stabilité.

Figure 2. Évolution temporelle de la densité d'états $\rho(\epsilon)$ d'un verre d'électrons modèle après une trempe à basse température. Les paramètres choisis correspondent à une interface MOS à base de In_2O_3 et les temps d'observation s'étalent entre le millionième de seconde et un peu plus d'un jour. Courbe encadrée : dérive temporelle de la conductivité σ . La densité d'états et la conductivité sont représentées en unités arbitraires, les énergies en kelvins.



matière atome par atome au moyen d'une pointe dont l'extrémité, de dimension atomique, se déplace au-dessus de la surface à observer (encadré 3).

L'interprétation des informations fournies requiert toutefois une approche théorique solide. Tout l'intérêt du microscope à effet tunnel réside en effet dans sa sensibilité aux détails des structures géométrique et électronique de la surface, permettant ainsi d'en obtenir des images avec une résolution à l'échelle atomique. Cependant, la visualisation par STM ne permet pas toujours une interprétation sans ambiguïtés de la surface observée, car la structure atomique et électronique de la surface ne peut pas être directement déduite de l'image obtenue. C'est le rôle de la théorie et de son interaction avec l'expérience de permettre une identification et une interprétation correcte des informations STM par le biais de la **simulation numérique** (encadré A. *Qu'est-ce qu'une simulation numérique ?*). Les théoriciens utilisent pour ce faire des méthodes de calcul numérique qui combinent une approche classique pour la dynamique des atomes et une description quantique (encadré 2) des électrons participant aux liaisons chimiques entre les atomes. À partir d'une description initiale plausible de la surface (composition chimique et géométrie de départ), ces méthodes permettent d'obtenir la structure géométrique d'équilibre du système ainsi que la configuration électronique correspondante. Il s'agit ensuite de calculer le courant circulant entre la pointe et la surface. Les formalismes théoriques (plus ou moins élaborés) développés pour le transport d'électrons dans les structures microscopiques sont alors utilisés et adaptés au problème du transport électronique dans la jonction STM. Le calcul du courant passant de la pointe vers la surface (ou *vice versa* selon le signe de la tension appliquée) peut être effectué. Les images (à hauteur de pointe constante ou à courant constant) sont calculées et comparées aux images expérimentales. Le processus est réitéré jusqu'à obtenir un accord suffisant entre expérience et théorie.

Modéliser la lenteur d'un "verre d'électrons"

Le troisième exemple porte sur la modélisation des "verres d'électrons", dans le domaine de la conduction électrique au sein des composants microélectroniques. Cette dénomination est donnée à des systèmes électroniques dont la dynamique, très lente à l'échelle des temps expérimentaux, ressemble à celle des verres structuraux. Dans des systèmes tels que les couches minces métalliques ou les interfaces métal-oxyde-semiconducteur (MOS), les électrons de conduction sont confinés dans un plan par la géométrie du système. Or des mesures de transport ont montré qu'il existe des systèmes où les électrons ne forment pas un fluide conducteur (comme dans les métaux ordinaires) mais un verre d'électrons. Un tel verre montre d'autre part des effets de mémoire : ses propriétés à un instant donné ne dépendent pas seulement des conditions dans lesquelles il se trouve (température, densité, etc.) mais aussi du chemin suivi pour les atteindre à partir d'un état initial. Les systèmes où ces propriétés remarquables ont été observées ont des résistances très élevées (plusieurs millions d'ohms), signe d'une structure très désordonnée. En présence d'un très fort désordre, les électrons subissent de subtils effets d'interférence quantique qui les empêchent de diffuser dans le système, comme dans un cristal faiblement désordonné. À cause de ces effets, particulièrement importants dans les systèmes bidimensionnels, les électrons restent piégés autour de certaines positions

distribuées au hasard dans l'espace. Les énergies de ces états, dits *localisés*, sont, elles aussi, aléatoires et s'étalent sur un intervalle dont la largeur est déterminée par le degré de désordre. Chaque état global du système correspond à une des multiples façons de distribuer les électrons parmi les états localisés. Les configurations de basse énergie du système résultent de la recherche d'un compromis entre deux conditions qui peuvent s'avérer incompatibles. D'une part, il convient de placer les électrons le plus loin possible les uns des autres afin de minimiser leur répulsion **électrostatique**. D'autre part, ils doivent occuper les états localisés les plus profonds en énergie afin de minimiser leurs interactions avec le potentiel aléatoire. Lorsque ces deux interactions sont du même ordre, il devient impossible de satisfaire aux deux conditions à la fois. On dit alors que le système est *frustré*.

Une conséquence importante de la frustration est la prolifération d'états métastables⁽³⁾ dont les énergies sont très proches, mais dont les structures sont très différentes. À basse température, en évoluant vers l'équilibre à partir d'un état initial donné, le système doit explorer successivement ces états. Or, les barrières de potentiel qui les séparent (et, par conséquent, le temps nécessaire pour effectuer une transition entre deux de ces états) sont si largement distribuées que, même après

(3) L'état métastable est l'état dans lequel un système physique est "bloqué" dans un état excité (à un niveau d'énergie supérieur à son état fondamental) pendant un certain temps.



La technique du microscope à effet tunnel

3

La technique du microscope à effet tunnel (ou STM, pour *scanning tunneling microscope*) repose sur l'utilisation d'une pointe dont l'extrémité est de dimension atomique (de plusieurs dizaines d'atomes à un atome unique) et qui se déplace au-dessus de la surface avec une précision extrême (inférieure au centième de nanomètre). En appliquant une tension entre l'échantillon et la pointe, un courant électrique est extrait pour établir une cartographie de la surface (figure A). Les images STM contiennent toujours un mélange subtil d'informations sur la topologie de la surface et sa capacité à transporter localement plus ou moins bien le courant électrique (encadré 2). Cependant il n'existe pas de théorème d'inversion permettant de remonter, de manière univoque, d'une image STM à la structure atomique et électronique de la surface. D'où le recours obligé à la théorie pour l'interpréter. La figure B donne l'exemple particulier d'une image simulée pour une surface de silicium sur laquelle sont déposées des molécules d'éthylène (C_2H_4).

Récemment, la manipulation d'atomes ou de molécules sur une surface a été effectuée à l'aide d'une pointe STM. Dans ce cas, le microscope à effet tunnel représente l'"outil ultime" pour l'élaboration, de manière contrôlée et reproductible, de structures artificielles sur les surfaces. Il ouvre un champ d'investigation fructueux pour la compréhension des propriétés quantiques de ces petits objets. Les scientifiques envisagent dès lors des applications technologiques. Déjà, le développement d'une électronique à l'échelle atomique ou moléculaire, par exemple, est en plein essor. Dans cette électronique du futur, les fonctions logiques des composants ou, plus simplement, le transport du courant seront assurés

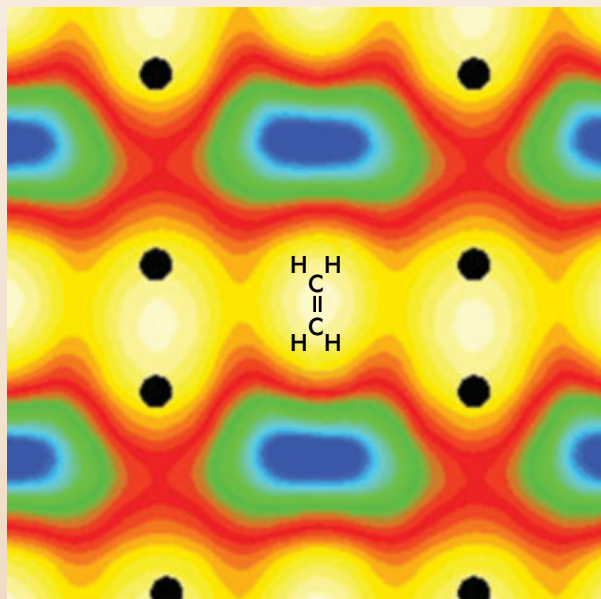


Figure B. Simulation de l'image à courant constant d'une surface de silicium (Si) sur laquelle repose une molécule d'éthylène (C_2H_4). Dans le cas présent, même si les molécules sont déposées sur les atomes de Si de la surface (points noirs), elles apparaissent moins brillantes que la surface nue, en accord avec les expériences.

par des molécules individuelles ou par des lignes d'atomes reposant sur des surfaces. L'étude de tels systèmes est en cours au sein du groupe MSIN.

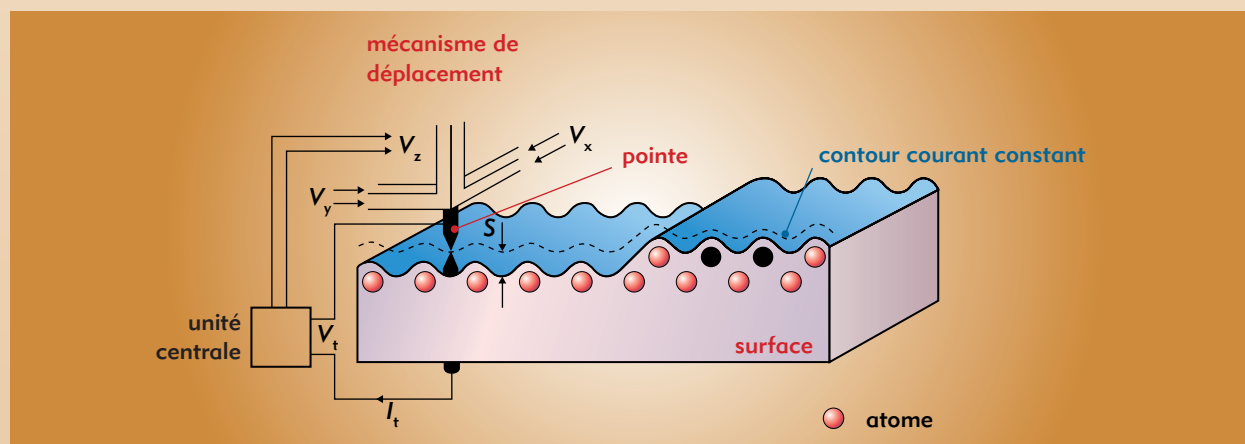


Figure A. Représentation schématisée du microscope à effet tunnel. Le mécanisme de déplacement de la pointe permet de balayer la surface (V_x, V_y) et de maintenir le courant constant I_t en ajustant la distance à la surface (V_z).

avoir attendu des temps arbitrairement longs, il reste toujours des barrières à traverser. Ce phénomène est à l'origine de la dynamique lente observée dans les expériences.

Cette lenteur est illustrée par l'évolution dans le temps de la densité d'états⁽⁴⁾ dans un modèle simplifié de verre d'électrons (figure 2). Le système se trouve à l'instant initial dans un état désordonné à très haute température à partir duquel il est trempé⁽⁵⁾ brutalement à une température de 4 K. On remarque l'extrême lenteur de l'évolution de la densité des états dont l'énergie ϵ est proche du niveau de Fermi μ .

En effet ceux-ci n'ont toujours pas atteint l'équilibre plus d'un jour après la trempé ! La conductivité du système, déterminée par ces états est très sensible à ce phénomène.

(4) La densité d'états est fonction de l'énergie $\rho(\epsilon)$ telle que $\rho(\epsilon) \Delta\epsilon$ est le nombre d'états électroniques d'un système dont l'énergie est comprise entre ϵ et $\epsilon + \Delta\epsilon$, où $\Delta\epsilon$ est un très petit intervalle d'énergie.

(5) L'opération de trempé consiste à figer, par un refroidissement brutal, une structure obtenue à haute température dans un alliage ou un verre.

Ceci se traduit par une dérive très lente dans le temps. Ici encore, la modélisation doit donc prendre en compte les caractéristiques de la matière avec une très grande précision. ●

Marie-Catherine Desjonquères,
Daniel Grepel, Hervé Ness
et Cyrille Barreteau
Direction des sciences de la matière
CEA centre de Saclay
Daniel Spanjaard
Laboratoire de physique des solides
Orsay

Qu'est-ce qu'une simulation numérique ?

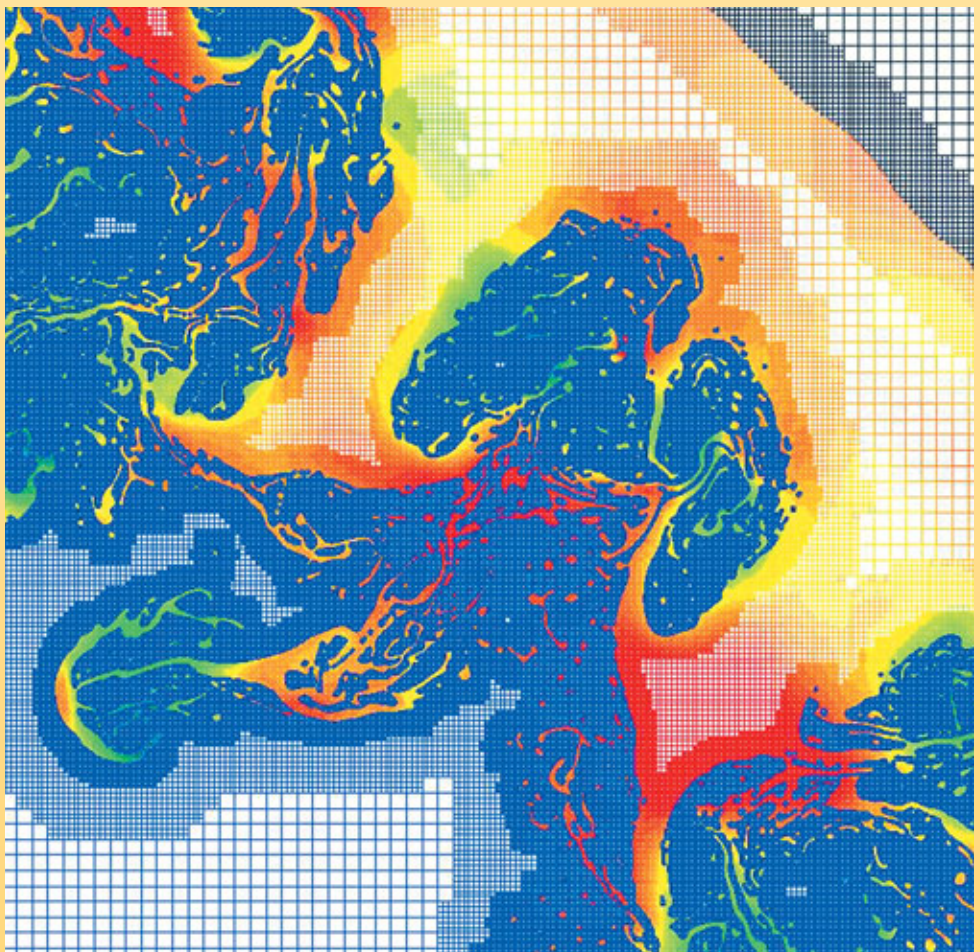
La **simulation numérique** consiste à reproduire par le calcul le fonctionnement d'un système, préalablement décrit par un ensemble de **modèles**. Elle s'appuie sur des méthodes mathématiques et informatiques spécifiques. Les principales étapes de la réalisation d'une étude par simulation numérique sont communes à de nombreux secteurs de la recherche et de l'industrie, en particulier le nucléaire, l'aérospatial ou l'automobile.

En chaque point de l'"objet" considéré, plusieurs grandeurs physiques (vitesse, température...) décrivent l'état et l'évolution du système étudié. Celles-ci ne sont pas indépendantes, mais reliées et régies par des **équations**, généralement **aux dérivées partielles**. Ces équations constituent la traduction mathématique des lois de la physique qui modélisent le comportement de l'objet. Simuler l'état de ce dernier, c'est déterminer – idéalement en tout point – les valeurs numériques de ses **paramètres**. Comme il y a un nombre infini de points, donc une infinité de valeurs à calculer, cet objectif est inaccessible (sauf dans des cas bien particuliers où l'on peut résoudre les équations de départ à l'aide de formules analytiques). Une approximation naturelle consiste donc à ne considérer qu'un nombre fini de points. Les valeurs des paramètres à calculer sont ainsi en nombre fini et les opérations nécessaires deviennent abordables grâce à l'ordinateur. Le nombre effectif de points traités dépendra bien sûr de la puissance de celui-ci : plus il sera élevé, meilleure sera finalement la description de l'objet. À la base du calcul des paramètres comme à la base de la simulation numérique, il y a donc la réduction de l'infini au fini, la **discrétisation**.

Comment opère-t-on précisément à partir des équations mathématiques du modèle ? Deux méthodes sont très souvent utilisées, respectivement représentatives des méthodes de **calcul déterministe**, qui résolvent les équations régissant les phénomènes étudiés après avoir discrétisé les variables, et des méthodes de **calcul statistique** ou **probabiliste**.

Le principe de la première, connue sous le nom de **méthode des volumes finis**, est antérieur à l'usage des ordinateurs. Chaque point de l'objet est assimilé simplement à un petit volume élémentaire (un cube par exemple), d'où le nom de *volume fini*. Un plasma, par exemple, est ainsi vu comme un ensemble ou un réseau de volumes contigus qui, par analogie avec la trame d'un tissu, sera dénommé **maillage**. Les paramètres de l'état de l'objet sont maintenant définis dans chaque maille du maillage. Pour chacune d'elles, en reformulant les équations mathématiques du modèle par des moyennes volumiques, il sera alors possible de construire des *relations algébriques* entre les paramètres de la maille et ceux de ses voisins. Au total, il y aura autant de relations que de paramètres inconnus et ce sera à l'ordinateur de résoudre le *système* de relations obtenu. Il faudra pour cela recourir aux techniques de l'**analyse numérique** et programmer des **algorithmes** spécifiques.

L'accroissement de la puissance des ordinateurs a permis d'augmenter la finesse de discrétisation, permettant de passer de quelques dizaines de mailles dans les années soixante à plusieurs dizaines de milliers dans les années quatre-vingt, à des millions dans les années quatre-vingt-dix et jusqu'à la dizaine de milliards de mailles aujourd'hui (machine Tera de la Direction



Exemple d'image d'une simulation 2D d'instabilités réalisée avec le supercalculateur Tera du CEA. Le calcul a fait appel au maillage adaptatif, qui se fait plus fin dans les zones où les phénomènes sont les plus complexes.

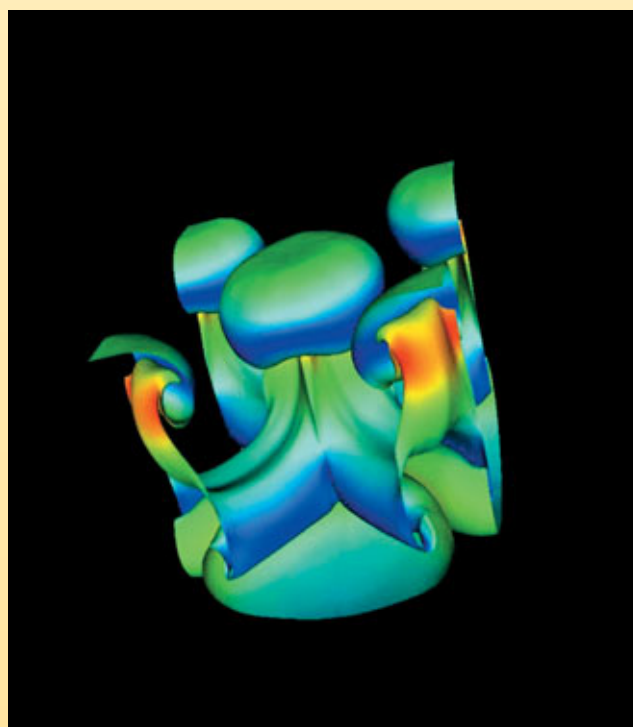
des applications militaires du CEA), chiffre qui devrait décroître à la fin de la décennie.

Un raffinement du maillage, le **remaillage adaptatif**, consiste à ajuster la taille des mailles en fonction des circonstances, par exemple en les rendant plus petites et plus serrées aux interfaces entre deux milieux, là où les phénomènes physiques sont les plus complexes, ou là où les variations sont les plus importantes.

La méthode des volumes finis s'applique dans des contextes physiques et mathématiques très variés. Elle autorise toute forme de maille (cube, hexaèdre, tétraèdre...) et le maillage peut être modifié durant le calcul, en fonction de critères géométriques ou physiques. Enfin, elle est aisée à mettre en œuvre dans le contexte des **ordinateurs parallèles** (encadré B, **Les moyens informatiques de la simulation numérique hautes performances**), le maillage pouvant en effet faire l'objet d'un découpage pour des calculs sur ce type de machines (exemple figure B, p. 13).

Appartiennent à la même famille la **méthode des différences finies**, cas particulier de la **méthode des volumes finis** où les côtés des mailles sont orthogonaux, et la **méthode aux éléments finis**, qui peut juxtaposer divers types de mailles. La deuxième grande méthode, dite de **Monte-Carlo**, est particulièrement adaptée pour simuler le *transport de particules*, par exemple des neutrons ou des photons d'un **plasma** (voir *Les simulations en physique des particules*). Un tel transport est en fait caractérisé par une succession d'étapes lors desquelles chaque particule peut subir différents événements (diffusion, absorption, émission...) possibles *a priori*. Les probabilités élémentaires de chacun de ces événements sont connues individuellement pour chaque particule.

Il est alors naturel d'assimiler un point du plasma à une particule. Un ensemble de particules, en nombre fini, va constituer un échantillon représentatif de l'infinité de particules du plasma, comme lors d'un sondage statistique. D'étape en étape, l'évolution de l'échantillon sera déterminée grâce à des tirages au hasard (d'où le nom de la méthode). L'efficacité de cette méthode, mise en œuvre à Los Alamos dès les années 1940, dépend bien sûr de la qualité statistique des tirages au hasard.



CEA

Simulation 3D réalisée à l'aide du supercalculateur Tera installé fin 2001 au centre CEA/DAM Île-de-France à Bruyères-le-Châtel (Essonne).

Il existe pour cela des méthodes de nombres aléatoires, bien adaptées au traitement par un ordinateur.

Les méthodes des volumes finis et de Monte-Carlo ont suscité et suscitent de nombreuses études mathématiques. Ces études s'attachent notamment à préciser la convergence de ces méthodes, c'est-à-dire comment la précision de l'approximation varie avec le nombre de mailles ou de particules. Cette question est naturelle lors de la confrontation des résultats de la simulation numérique à ceux de l'expérience.

Comment se déroule une simulation numérique ?

Il est souvent question d'*expérience numérique* pour souligner l'analogie entre la pratique d'une simulation numérique et la conduite d'une expérience de physique.

Brièvement, cette dernière utilise un dispositif expérimental, configuré selon des conditions initiales (de température, de pression...) et des paramètres de contrôle (durée de l'expérience, des mesures...). Durant l'expérience, le dispositif produit des points de mesures qui sont enregistrés. Ces enregistrements sont ensuite analysés et interprétés.

Dans une simulation numérique, le dispositif expérimental consiste en un ensemble de programmes informatiques exécutés sur des ordinateurs. Les **codes** ou **logiciels de calcul** sont la traduction, à travers des algorithmes numériques, des formulations mathématiques des modèles physiques étudiés. En amont et en aval du calcul, les *logiciels d'environnement* effectuent la gestion de plusieurs opérations complexes de préparation des calculs et de leur dépouillement.

Les données initiales de la simulation comporteront d'abord la délimitation du domaine de calcul à partir d'une représentation approchée des formes géométriques (produite par le dessin et la CAO, conception assistée par ordinateur), suivie de la discrétisation de ce

domaine de calcul sur un maillage, ainsi que les valeurs des paramètres physiques sur ce maillage et les paramètres de contrôle du bon déroulement des programmes... Toutes ces données (produites et gérées par les logiciels d'environnement) seront saisies et vérifiées par les codes. Les résultats des calculs proprement dits, c'est-à-dire les valeurs numériques des paramètres physiques, seront sauvegardés au fur et à mesure. En fait, un protocole spécifique structurera les informations produites par l'ordinateur afin de constituer une base de données numériques.

Un protocole complet organise l'échange informatique des informations requises (dimensions notamment) suivant des formats prédéfinis : *modeleur*⁽¹⁾, *mailleur*⁽²⁾, *découpeur de maillage*, *code*

(1) Le *modeleur* est un outil qui permet la création et la manipulation de points, courbes et surfaces en vue par exemple de la création d'un maillage.

(2) Les formes géométriques d'un maillage sont décrites par des ensembles de points reliés par des courbes et des surfaces (de Bézier par exemple) qui en représentent les frontières.



de calculs, logiciel de visualisation et d'analyse. Les études de *sensibilité* des résultats (au maillage et aux modèles) font partie des "expériences" numériques.

À l'issue des calculs (résolution numérique des équations décrivant les phénomènes physiques qui se déroulent dans chaque maille), l'analyse des résultats par des spécialistes reposera sur l'exploitation de la base de données numériques. Elle comportera plusieurs étapes : extraction sélective des données (selon le paramètre physique recherché) et visualisation, extraction et transfert des données pour calculer et visualiser des diagnostics.

Le parallèle entre la conduite d'un cas de calcul, d'une expérience numérique et la conduite d'une expérience physique ne s'arrête pas là : les résultats numériques seront comparés aux résultats expérimentaux. Cette analyse comparative, effectuée sur la base de critères quantitatifs standardisés, fera appel et à l'expérience

et à l'art de l'ingénieur, du physicien, du mathématicien. Elle débouchera sur de nouvelles améliorations des modèles physiques et des programmes informatiques de simulation.

Bruno Scheurer

Direction des applications militaires
CEA centre DAM-Ile de France

Frédéric Ducros et Ulrich Bieder

Direction de l'énergie nucléaire
CEA centre de Grenoble

L'exemple d'un calcul de thermohydraulique

La mise en œuvre d'un protocole de simulation numérique peut être illustrée par les travaux réalisés par l'équipe de développement du logiciel de calcul **thermohydraulique** Trio U. Ces travaux se sont déroulés dans le cadre d'une étude faite en collaboration avec l'Institut de radioprotection et de sûreté nucléaire (IRSN). L'objectif était d'obtenir des données très précises pour fournir à l'ingénieur les sollicitations en température à la paroi des composants d'un réacteur à eau sous pression dans le cas d'un accident grave impliquant une circulation naturelle turbulente de gaz chauds. Cette étude requiert la modélisation simultanée d'effets "système" à grande échelle et de phénomènes **turbulents** à petite échelle (encadré F, **Modélisation et simulation des écoulements turbulents**).

Elle débute par la définition du modèle de calcul global (figure A), suivie de la réalisation du modèle CAO et du maillage correspondant avec des logiciels du commerce (figure B). Les maillages de plus de cinq millions de mailles exigent l'utilisation de puissantes stations graphiques. Dans cet exemple, le maillage d'un générateur de vapeur (figures C et D) a été découpé pour répartir les calculs sur huit processeurs d'un ordinateur parallèle du CEA : chaque couleur symbolise une zone affectée à un processeur particulier. Les calculs, dont les conditions aux limites sont données par un calcul "système" (Icare-Cathare), produisent des résultats qu'il appartient aux spécialistes d'interpréter. En l'occurrence, les visualisations sur des stations graphiques des valeurs instantanées des champs de vitesse montrent l'impact d'un panache chaud sur la plaque tubulaire du

générateur de vapeur (coupe dans le champ de vitesses à gauche de la figure E) et la température instantanée dans la boîte à eau (à droite).

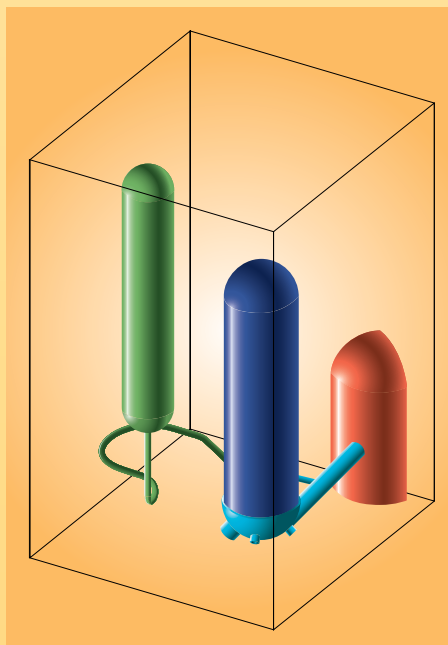


Figure A. Domaine de calcul global incluant une partie de la cuve réacteur (rouge), la conduite de sortie (branche chaude en bleu clair), le générateur de vapeur (bleu foncé) et le pressuriseur (vert).

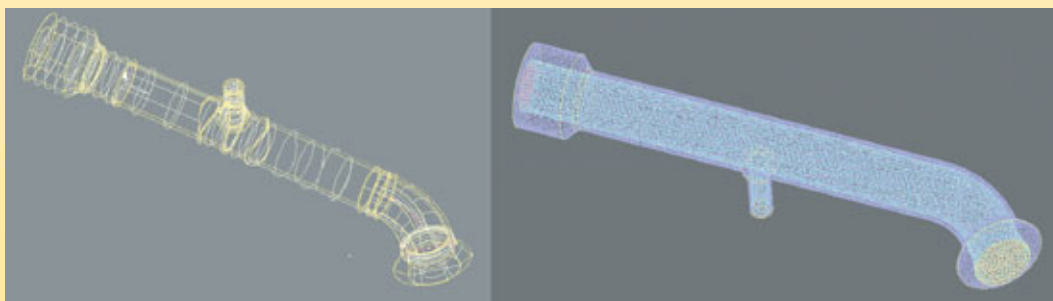
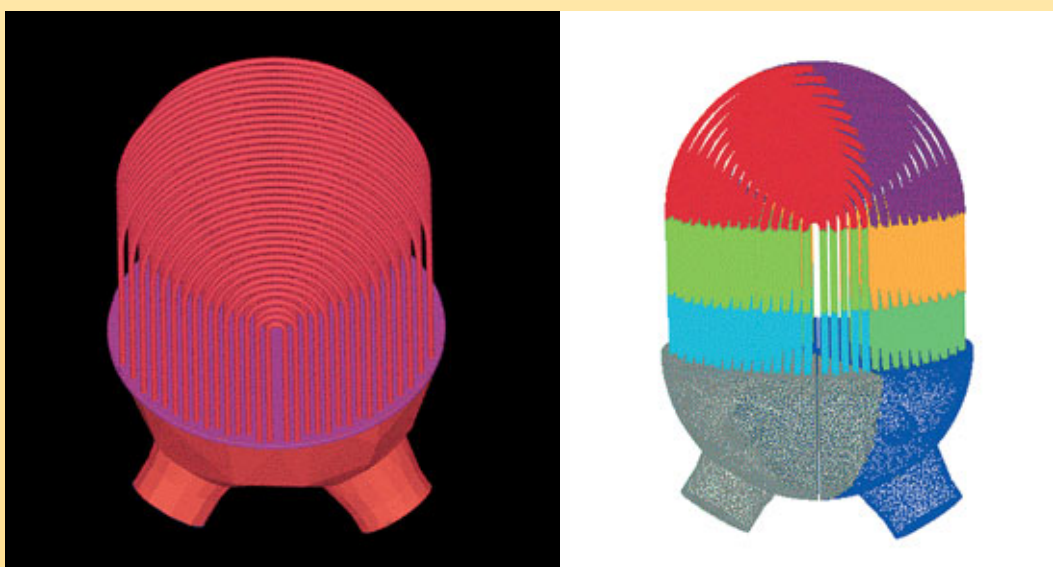


Figure B. Modèle CAO de la branche chaude en sortie de la cuve réacteur (à gauche) et son maillage non structuré (à droite).



Figures C et D.

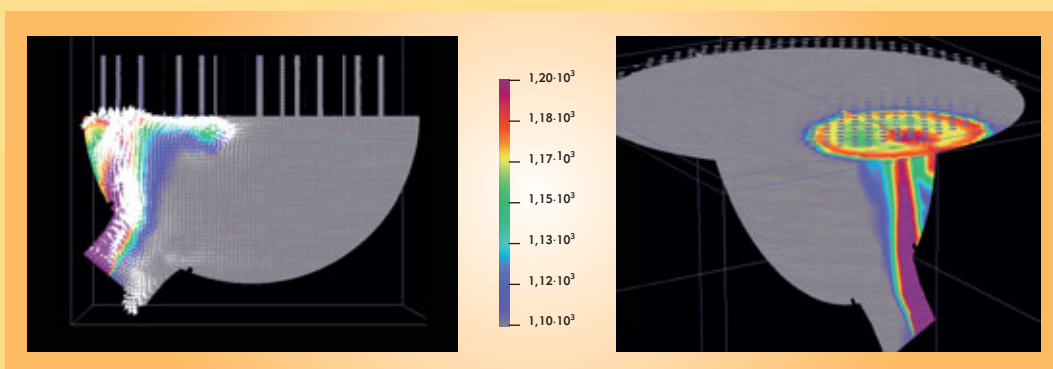


Figure E.

Les moyens informatiques de la simulation

Effectuer des **simulations numériques** plus précises impose de mettre en œuvre des **modèles** physiques et numériques eux-mêmes plus précis portant sur des descriptions plus fines des objets simulés (encadré A, *Qu'est-ce qu'une simulation numérique ?*). Tout ceci nécessite des progrès dans le domaine des logiciels de simulation mais aussi une augmentation importante de la capacité des équipements informatiques sur lesquels ces logiciels sont utilisés.

Processeurs scalaires et vectoriels

Au cœur de l'ordinateur, le processeur est l'unité de base qui, déroulant un programme, effectue les calculs. Il en existe deux grands types, les **processeurs scalaires** et les **processeurs vectoriels**. Les premiers exécutent des opérations portant sur des nombres élémentaires (scalaires), par exemple l'addition de deux nombres. Les seconds exécutent des opérations portant sur des ensembles de nombres (vecteurs), par exemple additionner deux à deux les nombres composant deux ensembles de 500 éléments. À ce titre, ils sont particulièrement adaptés à la simulation numérique : lors de l'exécution d'une opération de ce type, un processeur vectoriel peut fonctionner à une vitesse proche de sa performance maximale (crête). La même opération avec un processeur scalaire exige de nombreuses opérations indépendantes (opérations par composante des vecteurs) qui s'exécutent à une vitesse bien inférieure à sa vitesse crête. L'avantage principal des processeurs scalaires est leur prix : il s'agit de microprocesseurs généralistes dont les coûts de conception et de fabrication peuvent être amortis sur de larges marchés.

Forces et contraintes du parallélisme

Les processeurs récents permettent de hautes performances, d'une part en utilisant une fréquence de fonctionnement plus élevée, d'autre part en cherchant à exécuter en même temps

plusieurs opérations : c'est un premier niveau de **parallélisme**. L'accélération de la fréquence est limitée par l'évolution de la technologie micro-électronique, tandis que les dépendances entre instructions à exécuter par le processeur limitent le parallélisme possible. La mise en œuvre simultanée de plusieurs processeurs constitue un second niveau de parallélisme, qui permet d'obtenir des performances accrues à condition de disposer de programmes capables d'en tirer parti. Alors que le parallélisme au niveau des processeurs est automatique, celui *entre processeurs* dans un ordinateur parallèle est à la charge du programmeur, qui doit découper son programme en morceaux indépendants et prévoir entre eux les communications nécessaires. On procède souvent par un découpage du domaine sur lequel porte le calcul, chaque processeur étant chargé de simuler le comportement d'un domaine, et par l'établissement de communications régulières entre processeurs afin de garantir la cohérence d'ensemble du calcul. Pour obtenir un programme parallèle efficace, il faut s'assurer de l'équilibrage de charge entre processeurs et chercher à limiter le coût des communications.

Les différentes architectures

Les équipements informatiques ont différentes fonctions. À partir de son ordinateur de travail sur lequel il prépare ses calculs et en analyse les résultats, l'utilisateur accède à des moyens de calcul, de stockage, et de visualisation partagés, mais beaucoup plus puissants que les siens propres. L'ensemble de ces équipements sont reliés par des réseaux informatiques permettant de faire circuler les informations entre eux avec des débits compatibles avec le volume de données produites, pouvant atteindre 1 **téraoctet** (1 To = 10^{12} octets) de données pour une seule simulation.

Les grands équipements de calcul sont généralement appelés **supercalculateurs**. Ils atteignent aujourd'hui des puissances qui se chiffrent en **téraflops** (1 Tflops = 10^{12} opérations de calcul par seconde).

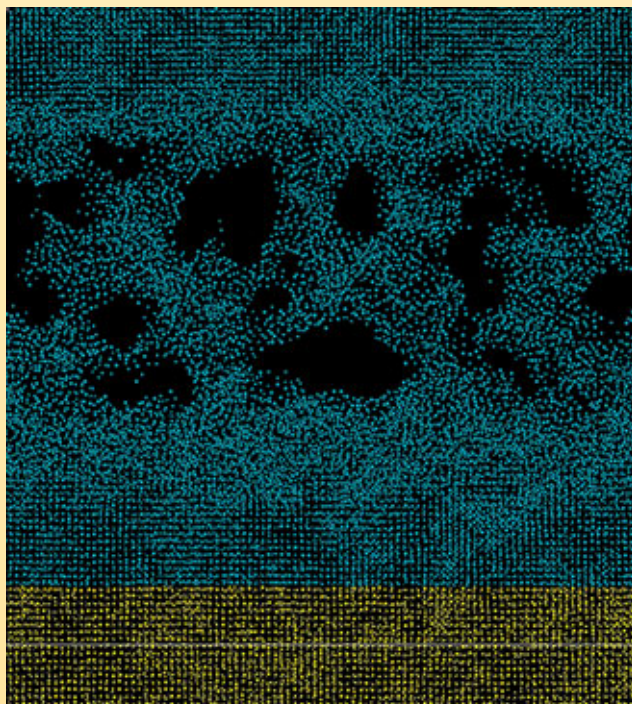
Il existe aujourd'hui trois grands types de supercalculateurs : les supercalculateurs vectoriels, les grappes de mini-ordinateurs à mémoire partagée et les grappes de PC (l'ordinateur que chacun possède chez soi). Le choix entre ces architectures dépend largement des applications et de l'utilisation visées. Les supercalculateurs vectoriels disposent de processeurs très performants mais dont il est difficile d'augmenter la puissance en ajoutant des processeurs. Les grappes de PC sont peu coûteuses mais mal adaptées à des environnements où de nombreux utilisateurs font beaucoup de calculs très gourmands en puissance machine, en mémoire et en entrées-sorties.

Ce sont ces considérations qui ont en particulier conduit la Direction des applications militaires (DAM) du CEA à choisir pour son programme simulation (voir *Le programme Simulation : la garantie des armes sans essais nucléaires*) les architectures de type grappe de mini-ordinateurs à mémoire partagée, encore appelées **clusters de SMP** (Symmetric Multi-Processor). Un tel système utilise comme brique de base un mini-ordinateur com-



CEA

Installée en décembre 2001 au CEA (centre DAM-Ile de France) et conçue par Compaq (devenue depuis HP), la machine Tera a pour élément de base un mini-ordinateur à 4 processeurs Alpha à 1 GHz partageant une mémoire de 4 Go et fournissant une puissance totale de 8 Gflops. Ces éléments de base sont interconnectés par un réseau rapide conçu par la société Quadrics. Une opération de synchronisation sur l'ensemble des 2 560 processeurs s'effectue en moins de 25 microsecondes. Le système de fichiers global offre un espace de stockage de 50 téraoctets pour les entrées-sorties avec une bande passante agrégée de 7,5 Go/s.



Les calculateurs parallèles sont adaptés aux méthodes numériques basées sur des maillages (encadré A, **Qu'est-ce qu'une simulation numérique ?**) mais aussi au traitement de calculs ab initio comme cette simulation par dynamique moléculaire de l'endommagement par choc de deux plaques de cuivre à 1 km/s (voir La simulation des matériaux). Le système considéré est constitué de 100 000 atomes de cuivre représentant un parallélépipède de section carrée (0,02 μm de côté) à densité normale. Les atomes interagissent suivant un potentiel EAM (embedded atom potential) pendant 4,6 picosecondes. Le calcul, effectué sur 18 processeurs du supercalculateur Tera de Bruyères-le-Châtel à l'aide du logiciel Stamp développé au CEA, a représenté une dizaine de minutes de temps "utilisateur" (calcul réalisé par B. Magne). Des tests impliquant jusqu'à 64 millions d'atomes ont été réalisés, mobilisant 256 processeurs pendant une centaine d'heures.

portant plusieurs microprocesseurs qui partagent une mémoire commune (figure). Ces mini-ordinateurs étant largement diffusés dans des domaines variés allant de la banque au serveur web

en passant par les bureaux d'études, ils offrent un excellent rapport performance/prix. Ces "briques" de base (encore appelées *nœuds*) sont reliées entre elles par un réseau d'interconnexion hautes performances : la puissance cumulée de plusieurs centaines de ces "briques" peut atteindre plusieurs téraflows. On parle alors d'**ordinateur massivement parallèle**.

Cette puissance peut être disponible pour une seule application parallèle utilisant toutes les ressources du supercalculateur mais aussi pour de multiples applications indépendantes, parallèles ou non, utilisant chacune une partie des ressources.

Si la caractéristique mise en avant pour décrire un supercalculateur est en général sa puissance de calcul, il ne faut pas négliger l'aspect entrées-sorties. Ces machines capables d'effectuer des simulations de grande taille doivent disposer de systèmes de disques avec des capacités et des performances adaptées. Dans les *clusters* de SMP, chaque mini-ordinateur dispose d'un espace disque local. Il n'est néanmoins pas judicieux d'utiliser celui-ci pour les fichiers utilisateurs, ce qui obligerait l'utilisateur à explicitement déplacer ses données entre les différentes phases de ses calculs. Pour cette raison, il est important de disposer d'un espace disque accessible par l'ensemble des mini-ordinateurs du supercalculateur. Cet espace est en général constitué de batteries de disques reliées à des nœuds dont la fonction principale est de les gérer. Comme pour le calcul, c'est le parallélisme pour les entrées-sorties qui permet d'offrir des performances élevées. Il faut, pour ce faire, disposer de systèmes de fichiers globaux parallèles permettant un accès rapide et sans contraintes à l'espace disque partagé.

Offrant des puissances de calcul considérables, les *clusters* de SMP posent néanmoins plusieurs défis. Parmi les plus importants, outre la programmation de logiciels de simulation capables de tirer parti du grand nombre de processeurs, il faut mettre au point des systèmes d'exploitation et les logiciels associés compatibles avec de telles configurations et tolérants vis-à-vis des pannes.

François Robin

Direction des applications militaires
CEA centre DAM-Ile de France

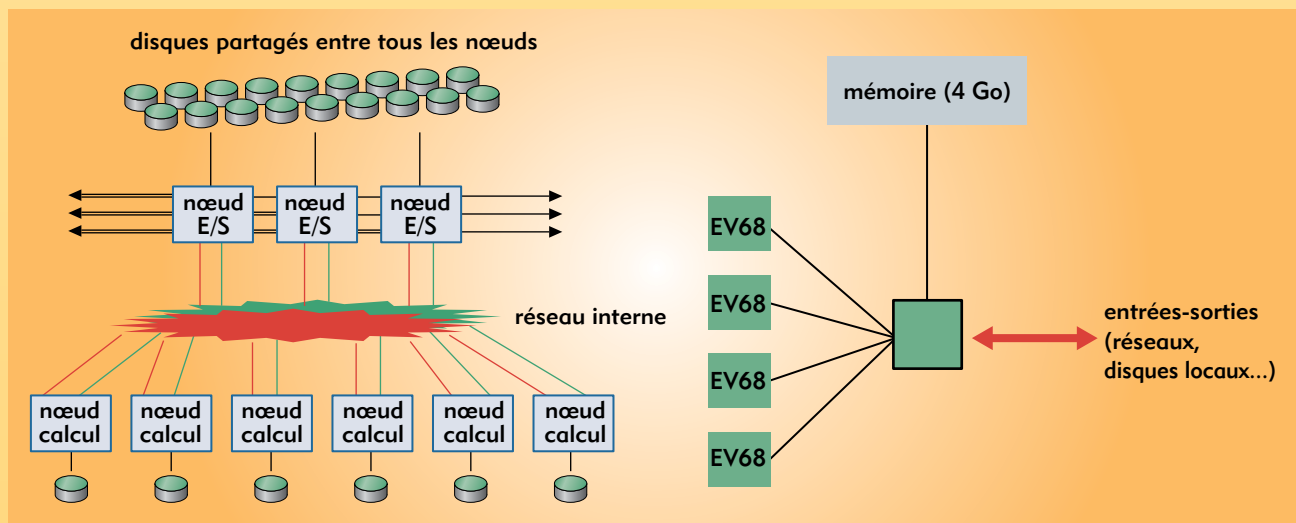


Figure. Architecture d'une machine du type "cluster de SMP". À gauche, l'architecture générale (E/S = entrée/sortie), à droite celle d'un nœud avec quatre processeurs Alpha EV68 cadencés à 1 GHz.

Modélisation et simulation des écoulements turbulents

La **turbulence**, ou l'agitation de l'écoulement dit turbulent, se développe dans la plupart des écoulements qui conditionnent notre environnement immédiat (rivières, océan, atmosphère). Elle se révèle être aussi un, sinon le, paramètre dimensionnant dans un bon nombre d'écoulements industriels (liés à la production ou la conversion d'énergie, à l'aérodynamique...). Il n'est donc pas étonnant que soient entrepris des efforts visant sa prédiction – fût-elle encore imprécise – surtout lorsqu'elle se trouve combinée à des phénomènes qui la compliquent : stratification, combustion, présence de plusieurs phases... C'est que, paradoxalement, même s'il est possible d'anticiper la nature turbulente d'un écoulement et même, d'un point de vue théorique, de dégager certaines caractéristiques communes et apparemment universelles aux écoulements turbulents⁽¹⁾, leur prédiction dans

des cas précis reste délicate. Celle-ci doit en effet prendre en compte l'importante gamme d'échelles spatiales et temporelles⁽²⁾ impliquées dans tout écoulement de ce type.

Les chercheurs ne sont pourtant pas démunis, aujourd'hui, pour aborder ce problème. En premier lieu, les équations qui régissent l'évolution spatio-temporelle des écoulements turbulents (équations de Navier-Stokes⁽³⁾) sont connues. Leur résolution complète, dans des cas très favorables, a conduit à des descriptions prédictives. Mais l'emploi systématique de cette méthode de résolution se heurte à deux difficultés rédhibitoires : d'une part, il nécessiterait la connaissance complète et simultanée de toutes les variables attachées à l'écoulement et des forçages s'exerçant sur lui⁽⁴⁾ et, d'autre part, il mobiliserait des moyens de calculs irréalistes pour encore des décennies.

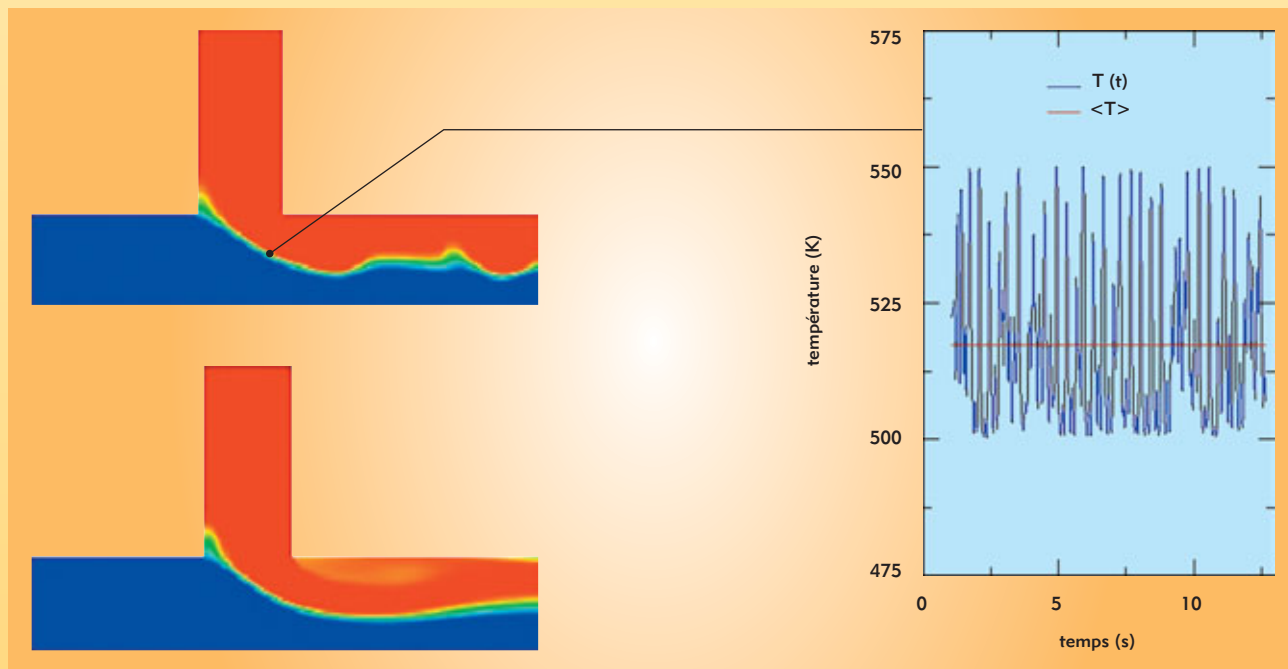


Figure. Champ de température instantané (haut) et moyenné (bas) dans une situation de mélange. La courbe donne l'historique de la température en un point : valeur instantanée fluctuante en bleu et moyenne en rouge (d'après la thèse d'Alexandre Chatelain [DEN/DTP/SMTH/LDTA]).

Il faut donc se résoudre, en s'appuyant sur le caractère fluctuant dû à l'agitation turbulente, à définir et utiliser des moyennes. Une des approches les plus répandues consiste à aborder le problème sous un angle statistique. Les moyennes d'ensemble de vitesse, de pression, de température... dont la distribution caractérise l'écoulement turbulent sont définies comme les variables principales de l'écoulement qu'on cherche à qualifier par rapport à ces moyennes. Ceci conduit à une décomposition du mouvement (dite de Reynolds) en champs moyen et fluctuant, ce dernier mesurant l'écart instantané et local entre chaque grandeur réelle et sa moyenne (figure). Ces fluctuations représentent la turbulence et couvrent une partie importante du spectre de Kolmogorov⁽¹⁾.

Cette opération réduit considérablement le nombre de degrés de liberté du problème et le rend « manipulable » informatiquement. Elle comporte aussi de nombreuses difficultés : il faut tout d'abord constater que, précisément en raison des non-linéarités des équations du mouvement, toute moyenne fait surgir des termes nouveaux et inconnus qu'il faut estimer. En fermant la porte à la description complète et déterministe du phénomène, on ouvre celle de la modélisation, c'est-à-dire à la représentation des effets de la turbulence sur les variables moyennes.

Beaucoup de progrès ont été accomplis depuis les premiers modèles (Prandtl, 1925). Les modélisations n'ont cessé d'évoluer vers plus de complexité, se basant sur le fait généralement vérifié que toute nouvelle extension permet de conserver les propriétés antérieurement acquises. Il faut aussi constater que, même

si de nombreux développements remettent en avant la nécessité de traiter les écoulements en respectant leur caractère *instationnaire*, les modélisations les plus populaires ont été développées dans le cadre des écoulements *stationnaires*, pour lesquels on n'accède donc qu'à une représentation de la moyenne temporelle de l'écoulement : dans le modèle mathématique final, les effets de la turbulence proviennent ainsi intégralement de la modélisation.

Il est également remarquable que, malgré de nombreux travaux, aucune modélisation n'est aujourd'hui capable de rendre compte de l'intégralité des phénomènes qui influencent la turbulence ou sont influencés par elle (transition, instationnarité, stratification, compression, etc.). Ce qui semble pour l'instant empêcher les modélisations statistiques de nourrir une ambition d'universalité.

Malgré ces limitations, la plupart des modélisations statistiques courantes sont maintenant disponibles dans les codes commerciaux et les outils des industriels. Il n'est pas possible de prétendre qu'elles permettent des calculs prédictifs dans toute situation. Leur précision est variable, offrant des résultats utiles pour l'ingénieur dans des situations maîtrisées et favorables (prédiction de la trainée avec une précision de 5 % à 10 % d'erreur [parfois mieux] sur certains profils), mais parfois faux dans des situations qui se révèlent, après coup, en dehors du champ de validité du modèle. Tout emploi maîtrisé d'une modélisation repose donc sur une qualification particulière au type d'écoulement à traiter. Des modélisations alternatives, répondant au besoin d'une plus grande précision sur des gammes d'échelles spatiales et temporelles plus étendues et donc basées sur un opérateur de "moyenne" d'une nature différente, sont actuellement en développement et représentent des voies nouvelles.

Le paysage des modélisations de la turbulence est aujourd'hui très complexe et l'unification des points de vue et des divers concepts de modélisation est une gageure. La tentation de l'universalité des modélisations reste donc hors de propos. Leur mise en œuvre réelle relève la plupart du temps de compromis généralement guidés par le savoir-faire de l'ingénieur.

Frédéric Ducros

Direction de l'énergie nucléaire
CEA centre de Grenoble

(1) On peut faire référence à la répartition spectrale de l'énergie cinétique turbulente, connue comme le "spectre de Kolmogorov", qui illustre de manière très simple la hiérarchie des échelles, des grandes échelles porteuses d'énergie aux échelles de plus en plus petites et de moins en moins énergétiques.

(2) Cette étendue est le résultat des non-linéarités des équations du mouvement qui donne naissance à une gamme étendue d'échelles spatiales et temporelles. Cette gamme est une fonction croissante du nombre de Reynolds, Re , mesurant le rapport entre force d'inertie et force visqueuse.

(3) L'hypothèse selon laquelle la résolution complète des équations de Navier-Stokes permet la simulation de la turbulence est généralement admise, tout du moins dans la gamme des écoulements sans choc.

(4) Il s'agit d'un problème régi par des conditions initiales et aux limites.