

Images de la physique 1972

Auteur(s) : CNRS

Les folios

En passant la souris sur une vignette, le titre de l'image apparaît.

81 Fichier(s)

Les mots clés

[Chabbal, Robert, 1927-2020, chambre à bulles](#)

Les relations du document

Collection Courrier du CNRS

Ce document est en lien avec :

[Le courrier du CNRS 7](#)

[Afficher la visualisation des relations de la notice.](#)

Citer cette page

CNRS, Images de la physique 1972, 1972

Valérie Burgos, Comité pour l'histoire du CNRS & Projet EMAN (UMR Thalim, CNRS-Sorbonne Nouvelle-ENS)

Consulté le 14/01/2026 sur la plate-forme EMAN :

<https://eman-archives.org/ComiteHistoireCNRS/items/show/253>

Copier

Présentation

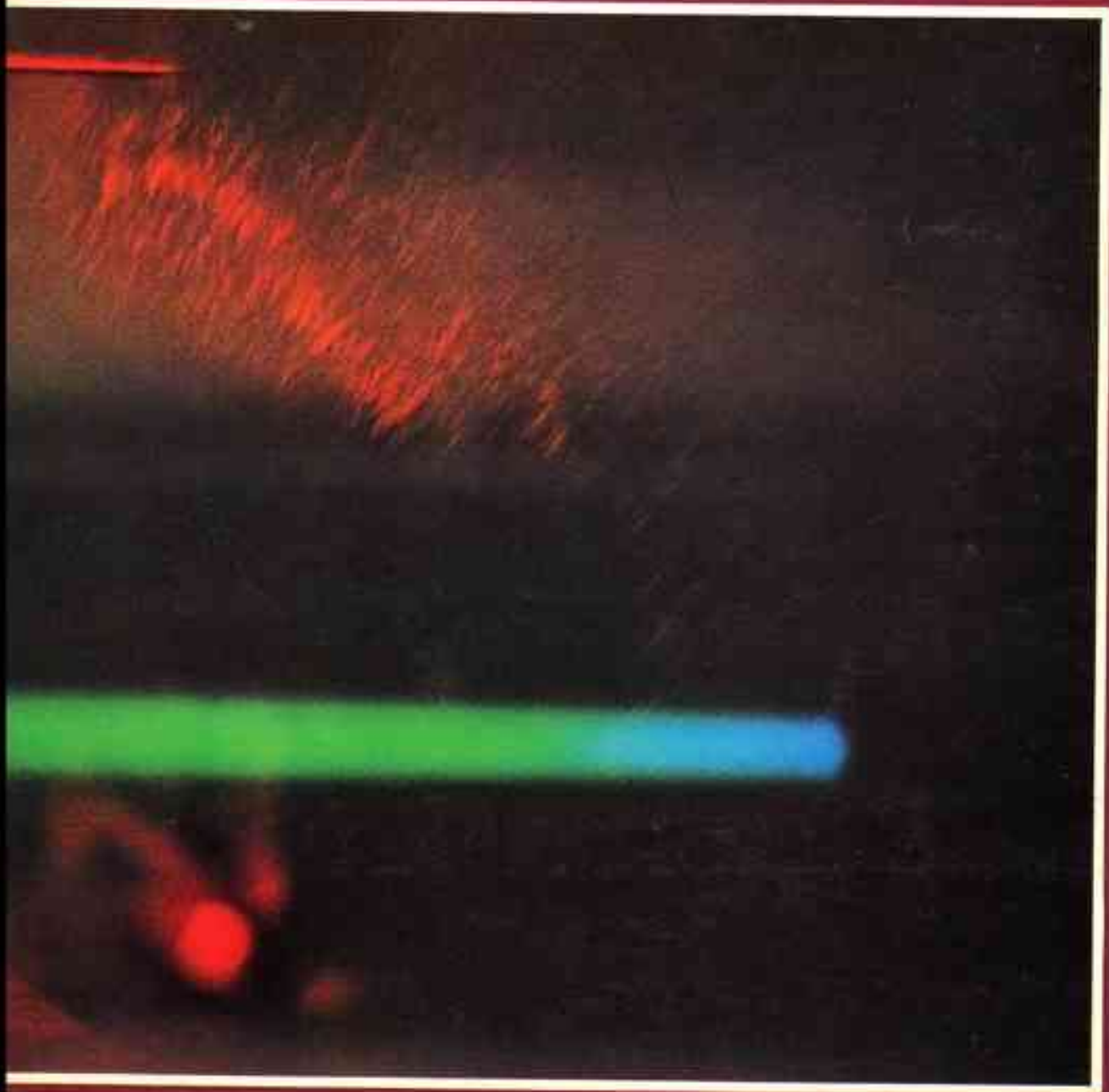
Date(s)1972

Mentions légalesFiche : Comité pour l'histoire du CNRS ; projet EMAN Thalim (CNRS-ENS-Sorbonne nouvelle). Licence Creative Commons Attribution - Partage à l'Identique 3.0 (CC BY-SA 3.0 FR).

Editeur de la ficheValérie Burgos, Comité pour l'histoire du CNRS & Projet EMAN (UMR Thalim, CNRS-Sorbonne Nouvelle-ENS)

Notice créée par [Valérie Burgos](#) Notice créée le 02/04/2024 Dernière modification le 24/12/2024

1972 IMAGES DE LA PHYSIQUE



SUPPLEMENT AU N° 7 DU COURRIER DU CNRS ^{CV}₂ PRIX 10 F

- 3 Présentation par Robert Chabbal, directeur scientifique du C.N.R.S.
- 4 Introduction

Noyaux et particules

- 5 La chambre à bulles "Gargamelle"
- 7 Faisceaux d'hyperons
- 9 Spectrométrie nucléaire à haute résolution avec les Protons de 1 GeV du synchrotron Saturne
- 12 Mise en évidence des noyaux critiques ou "Supermou" loin de la zone des noyaux stables

Atomes • Molécules • Matière condensée

- 15 Photodissociation d'ions H_2^+
- 17 Etude par coïncidence ion-photon des excitations dans les collisions $He^+ - He$
- 22 Spectroscopie "Beam Foil" et "Beam Gas"
- 25 Les lasers à haute stabilité de fréquence
- 27 Maser à hydrogène
- 31 Pompage optique des molécules
- 34 Pompage optique des électrons de conduction dans les semi-conducteurs
- 36 Pompage optique des centres colorés dans les cristaux
- 39 Condensation d'excitons en "gouttes" dans le Germanium
- 41 Anomalies de Kohn et distorsions de Peierls dans des conducteurs à une dimension
- 43 La diffraction des rayons X par les composés magnétiques
- 45 Structure des alliages métalliques amorphes
- 48 Diffusion Rayleigh inélastique et dynamique des petits mouvements dans les cristaux liquides
- 51 Dislocation de rotation dans les cristaux liquides
- 55 La vitesse critique dans l'Hélium super-fluide
- 58 Réfrigérateurs à dissolution et Physique des basses températures

Sciences de l'ingénieur

- 61 Les couches très minces de polymère
- 63 Modèles de composants électroniques semi-conducteurs
- 65 Identification par filtrage non linéaire
- 68 Commande des systèmes complexes
- 72 Conduite automatique du trafic urbain
- 75 Détection et prévention des pannes dans les circuits logiques

Supplément du numéro 7 du courrier du C.N.R.S.
avril 1973
Centre National de la Recherche Scientifique
15 quai Anatole-France - Paris 7^e - Tél. : 555.26.70
Directeur de la publication : René Audé
C.P. N° A.D.303 I.S. éditeur imprimé en France

photo de couverture

Un faisceau d'ions Li^+ d'énergie 1 MeV est excité par un mince film de carbone (à la droite de la photo) et se désexcite ensuite spontanément dans le vide. La lumière bleue émise est due aux ions Li^+ , la lumière verte aux ions Li^0 . La durée de vie correspondant à Li^0 étant plus longue, la lumière verte est observée sur un parcours plus long que la lumière bleue.

1972 IMAGES DE LA PHYSIQUE

La science connaît un développement rapide et il est nécessaire, mais très difficile de faire connaître aux non spécialistes l'extraordinaire richesse des phénomènes découverts. Le rapport d'activité du C.N.R.S. présente traditionnellement aux spécialistes une étude exhaustive des recherches en cours mais sa lecture est difficile : un groupe de scientifiques a donc bien voulu se charger d'en extraire quelques exemples significatifs et de les mettre à la portée de tous ceux qui s'intéressent aux questions scientifiques. Ce travail montre à quel point la Physique est riche en découvertes et en effets nouveaux et quelle place importante les laboratoires français tiennent dans le développement scientifique international.

Je souhaite que cette brochure contribue à démontrer que la Physique en 1972 est une passionnante aventure, et qui se passe parmi nous.

Robert Chabbal

La Physique est une vaste discipline et les recherches qui s'y mènent dans le cadre du CNRS sont aussi nombreuses que variées. Cette brochure ne prétend en aucun cas présenter un panorama complet de toutes les activités du CNRS dans cette matière, tâche à laquelle un gros volume ne suffirait pas, à moins de limiter la description de chaque travail de recherche à quelques mots peu intelligibles pour des non-spécialistes. C'est un point de vue diamétralement opposé qu'a adopté la petite équipe de journalistes et physiciens responsable de l'élaboration de ce document : un nombre relativement restreint de thèmes de recherches ont été retenus, de façon à pouvoir consacrer une place plus importante à chacun d'entre eux, ce qui permet de situer leur intérêt dans un cadre assez général. Le but poursuivi a donc été, en réunissant dans une brochure quelques expériences appartenant à des domaines variés, d'essayer de refléter pour un public assez large la diversité et l'intérêt des recherches en Physique effectuées soit dans le cadre du CNRS, soit avec son soutien. Dans cette optique, il a évidemment fallu faire un choix des sujets à retenir. Bien que toutes les expériences décrites concernent des travaux de pointe, ce choix n'a pas été uniquement guidé par des critères scientifiques ; on s'est également efforcé de tenir compte de la facilité avec laquelle semblaient pouvoir s'expliquer à des non-spécialistes tel ou tel sujet de recherches, son intérêt, les dispositifs expérimentaux mis en œuvre, etc... Il faut d'ailleurs reconnaître qu'un tel choix comporte toujours une grande part d'arbitraire. On voit à quel point cette brochure est éloignée d'un rapport d'activité exhaustif concernant l'ensemble

des recherches effectuées dans le cadre ou avec l'aide du CNRS : il s'agit encore moins d'un quelconque palmarès qui aurait la prétention de faire ressortir la "crème" des recherches en cours. Le lecteur remarquera d'ailleurs sans peine que des pans entiers de la Physique ne sont pas représentés dans ce document : la Physique Théorique, les Plasmas, l'Astrophysique, etc... Le renouvellement chaque année des physiciens de l'équipe de rédaction permettra d'obtenir, au bout de quelques années, une vue d'ensemble de la Physique qui se fait au CNRS. Pour cette année, trois grands thèmes ont été abordés :

"noyaux et particules"
 "atomes, molécules, matière condensée"
 "sciences pour l'ingénieur".

Dans la plupart des cas ces travaux ont été menés en collaboration par plusieurs chercheurs, ou même plusieurs laboratoires. Il était difficile de donner une liste complète de ceux qui avaient contribué à ces études ; nous avons donc pris le parti de ne citer que les noms des laboratoires concernés.

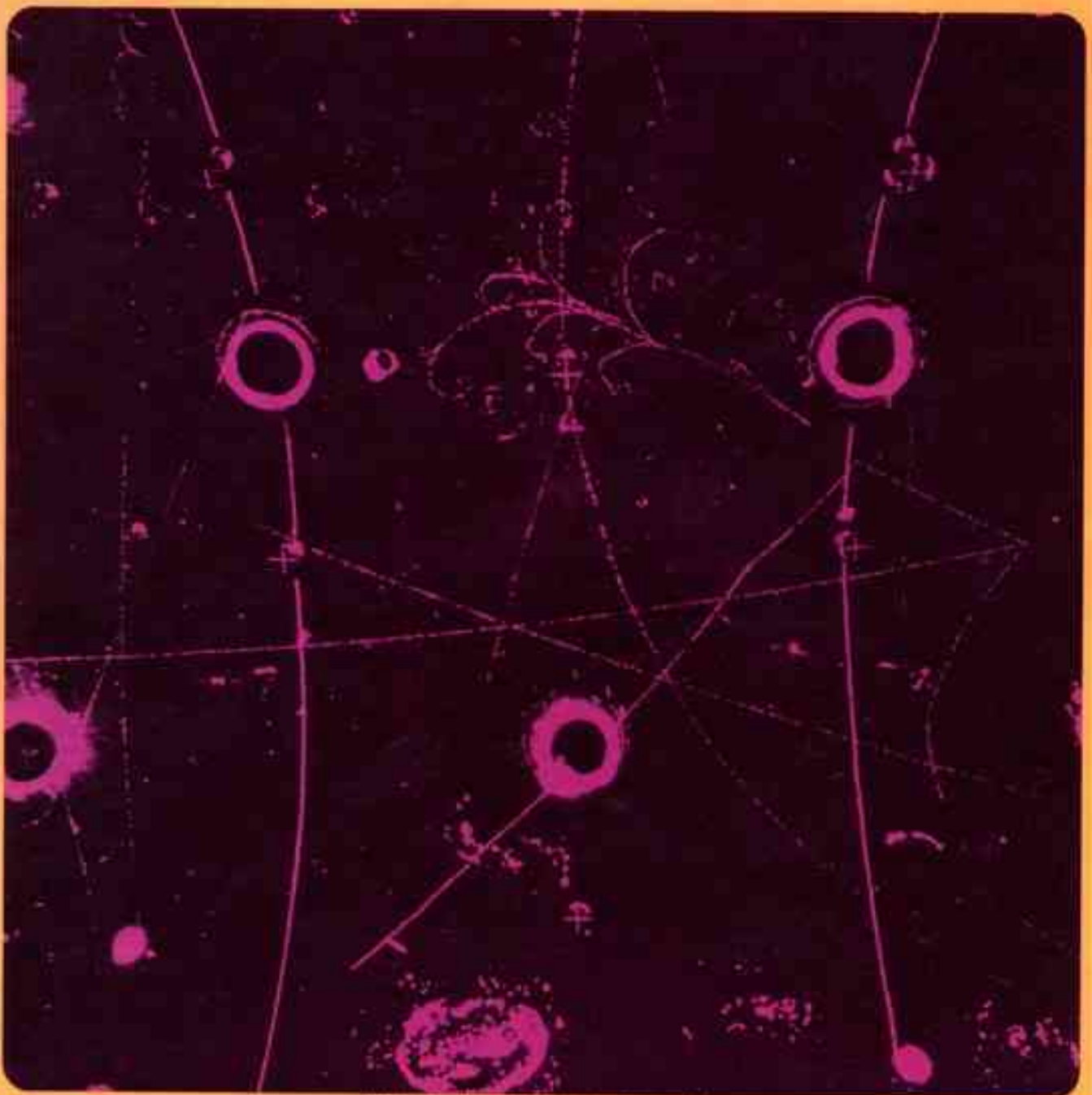
On trouvera ci-dessous la liste de ceux qui nous ont directement aidés, soit en rédigeant des textes, soit en nous fournissant des documents ; seule leur aide nous a permis de venir à bout de ce travail. Nous avons bénéficié, par l'intermédiaire de Madame M. Fitreman, de l'appui constant du Département des Relations Extérieures du CNRS. Nous remercions également Mademoiselle Dominique Vergèse, journaliste scientifique, pour ses précieux conseils.

Jean-Pierre Babary
 Elisabeth Dubois-Violette
 Robert Klapisch
 Franck Laloë
 Suzy Mouchet
 François Roubellat

ont apporté leur concours à l'équipe de rédaction :

J. Aguilar-Martin	J. Dixmier	M. Kléman	G. Rey
C. Audoin	M. Dufay	G. Lampel	A. Rousset
P. Azema	G. Durand	P. de Larminat	J. Saudinos
J. Baudon	J. Durup	J.-C. Lehmann	J. Thirion
F. de Bergevin	R. Foucher	A. Libchaber	D. Thoulouze
H. Carchano	P. G. de Gennes	H. Martinot	A. Tittl
B. Cheneveaux	A. Giraud	Y. Merle d'Aubigné	E. Varoquaux
R. Comes	F. Hartmann	J.-P. Repellin	M. Voos

noyaux et particules



la chambre à bulles "gargamelle"

Depuis que les accélérateurs de grande énergie produisent de nombreuses particules élémentaires, les chambres à bulles ont été des détecteurs particulièrement adaptés à l'étude du comportement de ces particules. C'est ainsi que les possibilités du synchrotron à protons de 25 GeV du CERN ont conduit à la construction de la chambre à bulles Gargamelle qui a fonctionné pour la première fois le 8 décembre 1970. La première expérience, commencée en mai 1971, a été terminée en novembre de la même année. La mise en évidence, pour la première fois, de la production par des antineutrinos de particules étranges isolées, a été publiée en juillet 1972 : la première détermination des sections efficaces des antineutrinos a été annoncée à la Conférence de Brookhaven en juin 1972 et à celle de Batavia en septembre 1972. Ces deux résultats fondamentaux viennent récompenser très largement les efforts du département Saclay de Saclay qui a construit la chambre, du groupe TC. L. du CERN qui a réalisé l'expérience, et de 7 laboratoires européens qui l'ont analysée. Dès 1963-1964, des physiciens de l'Ecole Polytechnique avaient proposé de construire une chambre à bulles à liquides lourds d'un volume utile qui soit de l'ordre de 10 fois plus grand que celui de la chambre précédente et ce en vue d'étudier les interactions de neutrinos. En effet, si les neutrinos sont produits en abondance auprès de l'accélérateur à 25 GeV, les sections efficaces sont si faibles qu'il fallait prendre, en 1963, environ un millier de clichés de chambre à bulles pour avoir une interaction. En 1971, on a obtenu dans Gargamelle une interaction pour 25 clichés.

Une chambre à bulles est une enceinte remplie de liquide dans des conditions assez proches de son point critique et maintenue à une pression supérieure à la pression de vapeur saturante. Un peu avant que les particules du faisceau ne pénètrent dans la chambre on détend brusquement celle-ci à une pression inférieure à la pression de vapeur saturante. Le liquide est alors dans un état métastable (retard à l'ébullition). Les particules chargées ionisent les molécules du liquide, provoquant des paires d'ions-électrons. Certains de ces électrons sont eux-mêmes très ionisants et, en fin de parcours, libèrent une énergie spécifique importante. Il s'ensuit un "point chaud", élément initial de la formation d'une bulle que l'on laisse

se développer jusqu'à une taille où elle est photographiable. On l'éclaire par un flash et le chapelet de bulles créé le long de la trajectoire de la particule permet de mesurer les caractéristiques cinématiques de celle-ci. Pour en mesurer la quantité de mouvement p , la chambre est dans un champ magnétique B intense, et la mesure de la courbure $1/R$ donne $p = RB$.

Les différents éléments de Gargamelle sont : un électro-aimant de 1000 tonnes d'une puissance égale à 6 MégaWatts ; un corps de chambre de 40 tonnes d'un volume intérieur de 12 m³ résistant à des pressions de 25 bars ; un système de détente avec un compresseur de 2,5 MégaWatts et un système d'éclairage constitué de 21 tubes flashs ; 8 caméras enfin fonctionnent toutes les deux secondes.

Gargamelle a été installée au CERN en 1969. Le corps de chambre est arrivé en juillet 1970 et la chambre a été détentée pour la première fois 5 mois après, donnant les premiers clichés des trajectoires de rayons cosmiques.

La première expérience

Pour sa première expérience, Gargamelle a fonctionné avec du fréon lourd (C₂F₅Br) permettant une densité relativement élevée (1,5 g/cm³). 500 000 clichés ont été pris, la moitié avec des neutrinos, l'autre moitié avec des antineutrinos.

Les interactions des neutrinos ont été très peu étudiées jusqu'à présent, si bien que cette expérience est intéressante d'abord par les informations aussi simples que les comparaisons de section efficace de neutrinos et d'antineutrinos, leur comportement en fonction de l'énergie, en fonction du transfert de moment et d'énergie au nucléon cible.

Si les interactions de neutrinos (ν) avec des protons (p) sont rares, les interactions d'antineutrinos ($\bar{\nu}$) le sont encore davantage (10 fois moins environ). De plus, avec les antineutrinos, on est particulièrement intéressé par la mise en évidence pour la première fois de la création d'une seule particule

étrange. Sur un premier lot de 1000 événements, 12 ont pu être identifiés comme production d'un hyperon Λ^0 isolé suivant la réaction :



Une telle réaction est tout simplement l'inverse de la désintégration muonique de Λ^0 . C'est une interaction faible, elle viole la règle de sélection selon laquelle « l'étrangeté » est conservée dans les réactions régies par les interactions fortes ». C'est pour cela qu'elle permet de produire un Λ^0 tout seul, qui n'est pas accompagné par une autre particule étrange, d'étrangeté opposée, un K^0 par exemple*.

La section efficace des antineutrinos sur les nucléons a été trouvée plus faible que celle des neutrinos. Le rapport entre les deux est compatible avec une valeur 1/3. Ce résultat préliminaire demande à être précisé, mais déjà il est très intéressant. En effet, les neutrinos et les antineutrinos de haute énergie sont des sondes très fines pour explorer le nucléon, et permettent de tester différents modèles de structure du nucléon que l'on suppose construit à partir de particules "plus" élémentaires, les partons. En particulier, le rapport 1/3 entre les sections efficaces d'antineutrinos et de neutrinos indique que ces partons doivent avoir un spin demi entier, c'est-à-dire que ce sont des fermions, et que la proportion des antipartons dans le nucléon doit être faible.

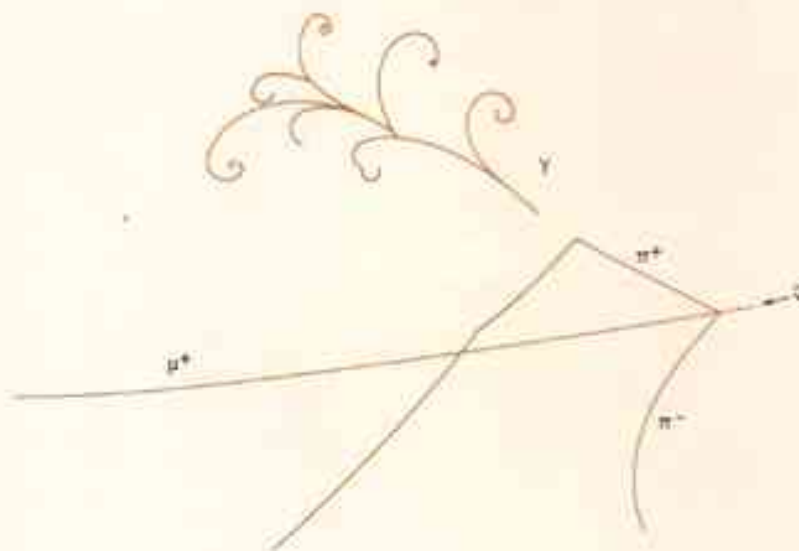
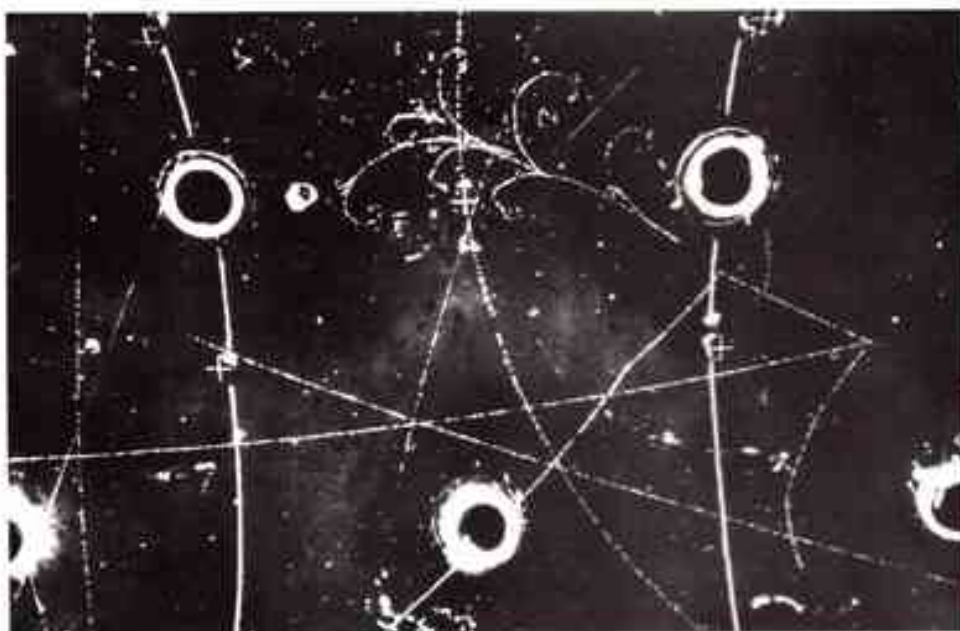
L'étude des clichés de neutrinos est en cours, elle permettra une analyse plus fine car les statistiques sont plus élevées. Une deuxième expérience faite avec Gargamelle remplie d'un mélange de propane et de fréon met en évidence l'importance de l'utilisation d'un liquide dense. Il s'agit de l'étude des annihilations d'antiprotons ($\bar{p}$) à 1,6 GeV/c et à 2,6 GeV/c avec la production de plusieurs pions neutres (π^0) et c'est le cas par exemple de la réaction :



Une telle réaction ne peut pas être étudiée en chambre à bulles à hydrogène. Le π^0 se désintègre en effet en deux photons, lesquels peuvent être détectés à haute énergie par leur matérialisation en paires d'électrons ($e^+ e^-$) se

* Département de Synchrotron Saclay (CERN Saclay) ;
Laboratoire de Physique des Particules Fondamentales de l'Ecole Polytechnique ;
Groupe TC. L. du CERN ;
Laboratoire de l'Accélérateur Linéaire - Orsay - Paris Sud.

* Les particules étranges sont toujours produites en paires dans les interactions fortes, par exemple $\pi^+ + p \rightarrow \Lambda^0 + K^+$. Le Λ^0 et le K^+ sont deux particules étranges avec des nombres quantiques d'étrangeté opposés.



Interaction d'un antineutrino et d'un proton dans Gargamelle.
Le muon positif est émis vers l'avant à grande énergie. Un π^+ et un π^- sont émis dans l'interaction. Le π^- s'arrête.
Le π^- réinteragit en produisant un π^0 dont on peut voir les γ matérialisés dans le liquide.

produisant au voisinage des noyaux lourds. Dans le liquide lourd de Gargamelle, les six photons donnent six paires d'électrons ($e^+ e^-$) et la réaction peut être identifiée et analysée. Dans l'hydrogène liquide, ces mêmes photons auraient échappé à la détection.

Perspectives d'avenir

Gargamelle est maintenant en bonne voie de réaliser le programme d'expériences qui a motivé sa construction. Ce programme se poursuivra jusqu'en 1974 sur l'accélérateur de 25 milliards d'électrons volts (25 GeV).

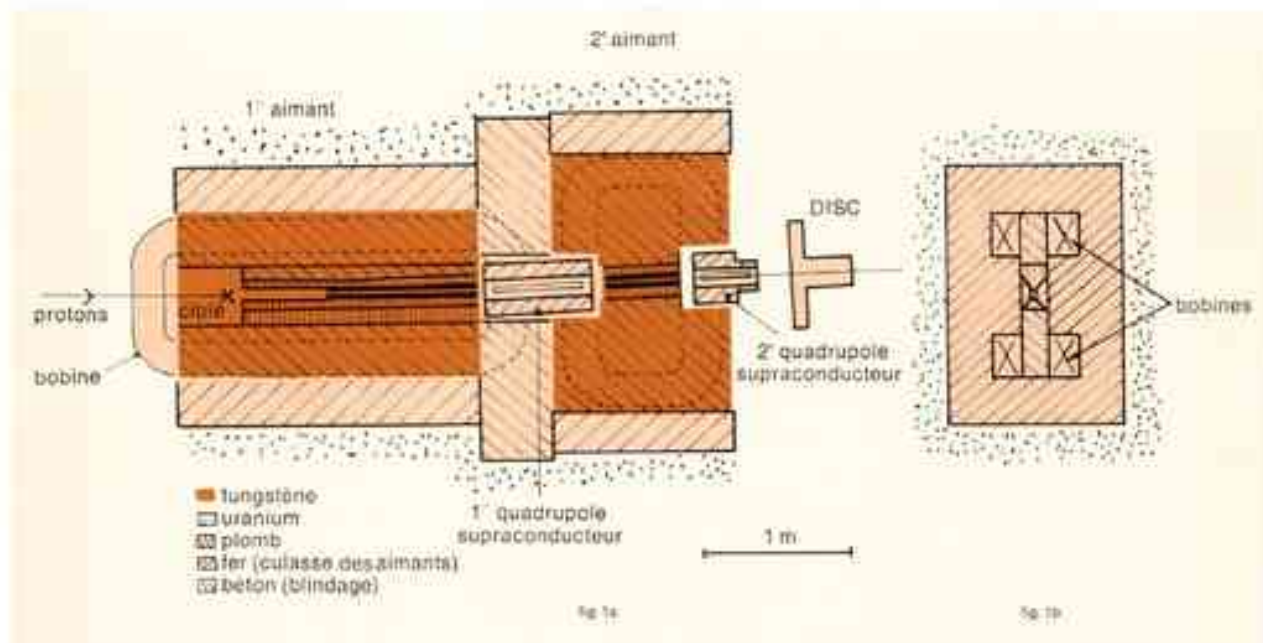
Mais en 1976, le nouvel accélérateur européen doit fournir des protons de 200 GeV. Ces protons peuvent être la source de faisceaux neutrinos de plus en plus intenses et d'énergie de plus en plus élevée. La section efficace des neutrinos augmentant linéairement avec l'énergie et, par ailleurs, la production des mésons π et K parents des neutrinos augmentant aussi avec l'énergie, le nombre d'interactions de neutrinos peut augmenter considérablement et atteindre une interaction par cliché

Il faut noter en passant que la forme allongée de la chambre est bien adaptée

pour une utilisation à très haute énergie. Il est donc très indiqué de prévoir l'installation de Gargamelle en ligne avec la grande chambre à hydrogène BEBC (Big European Bubble Chamber), actuellement en fin de construction. Elle pourra alors continuer essentiellement la physique des neutrinos à des énergies encore plus élevées.

Ces expériences ouvriront la voie à de nombreux domaines de physique inexplorés, par exemple la recherche du boson intermédiaire (responsable des interactions faibles) et la mise en évidence de l'interaction faible purement leptonique ($\bar{\nu}_\mu + e^- \rightarrow \mu^- + \nu_e$).

faisceaux d'hypérons



Le canal magnétique pour hypérons employé au CERN concentre de façon extrêmement compacte la plupart des éléments optiques d'un faisceau de particules à vie longue.

Les deux aimants donnent des champs magnétiques perpendiculaires au plan du dessin. On a représenté en couleur le contour des zones dans lesquelles le champ s'exerce. Ces champs magnétiques produisent une déflexion (visible sur le schéma) des particules comprises dans une certaine bande d'impulsion. Le rôle des deux quadrupôles est de rendre le faisceau de particules ainsi sélectionnées parallèle à l'axe. L'emploi de quadrupôles magnétiques supraconducteurs permet d'obtenir des éléments très compacts et avec un gradient de champ très élevé. Le compteur DISC permet la sélection des vitesses et à impulsion donnée on a donc réalisé une séparation en masse des différentes particules.

En 1b, on a représenté une section perpendiculaire du canal dans le premier aimant.

Les expériences de physique des hautes énergies n'utilisent pas directement en général les faisceaux de protons sortant des accélérateurs. On fait plutôt usage des particules "secondaires" créées par les interactions de protons sur une cible. Ces particules étant émises vers l'avant dans un cône angulaire étroit, un ensemble d'éléments optiques (électro-aimants, quadrupôles etc.) permet de les séparer (en quantité de mouvement, en charge), de les relocaliser etc. On obtient ainsi des "faisceaux secondaires" intenses de π , de K, voire d'anti-protons. Ces faisceaux s'étendent souvent sur de longues distances entre la cible de production et la cible d'expérience.

Les hypérons Σ , Ξ , Ω sont des particules dont la durée de vie est très courte et leur parcours moyen n'est que de quelques dizaines de centimètres au maximum. Il est donc difficile de construire des faisceaux secondaires de ces particules dans des conditions acceptables de distance à la cible et d'intensité. On peut surmonter ce handicap en créant des hypérons à aussi haute énergie que possible (pour profiter de la dilatation du temps) et en réalisant un faisceau aussi court que possible. A titre d'exemple, à haute impulsion, le parcours moyen des hypérons Σ et Ξ augmente à peu près propor-

tionnellement à l'impulsion et vaut environ 80 cm à 20 GeV/c. Quant aux Ω , leur parcours à cette énergie est de l'ordre de 40 cm. L'existence de faisceaux de protons d'énergie supérieure à 24 GeV a permis depuis trois ans la construction de faisceaux secondaires d'hypérons à Brookhaven et au CERN. Le faisceau réalisé au CERN utilise un canal magnétique constitué de deux aimants et de deux quadrupôles supraconducteurs (figure 1). La longueur totale de ce canal est de 360 cm et peut produire des hypérons chargés négativement, Σ^- , Ξ^- et éventuellement Ω^- jusqu'à une énergie de 20 GeV. Le canal est entouré de blindages de densités décroissantes (uranium et tungstène au centre, plomb et fer à la périphérie).

Ce canal assure une première sélection en moment des particules. Mais les hypérons sont accompagnés d'un flux très important d'autres particules, des pions essentiellement. Il est de ce fait très utile de pouvoir "signer" sélectivement les hypérons. Ceci est réalisé à l'aide d'un compteur Cerenkov diffé-

rentiel (DISC) dont le principe de fonctionnement est basé sur la sélection de vitesse. L'acceptance de ce compteur est limitée angulairement aux particules qui ne s'écartent pas trop de l'axe du compteur. Pour cette raison, le canal magnétique comprend deux lentilles quadrupolaires qui donnent un faisceau parallèle à la sortie du canal. Le compteur DISC a une longueur totale de 40 cm, et une section utile de détection de 3 cm de diamètre. En faisant varier la pression de l'hexafluorure de soufre qui remplit le DISC, il est possible de varier la bande des vitesses acceptées par le compteur et donc - à moment constant - de sélectionner successivement toutes les particules du faisceau, des pions jusqu'aux Ω^- .

Mesure des sections efficaces

Dans une première expérience, on a mesuré les sections efficaces de production des K^+ , $\bar{p}$, Σ^+ et Ξ^+ par comparaison avec celle des π^+ produits par des protons de 24 GeV sur une cible de tungstène. Ces mesures ont été faites à un angle de production fixe de 10 milliradians et des impulsions comprises entre 13 et 20 GeV/c, puis à une impulsion fixée à 17 GeV/c et des angles variant entre 10 et 40 milliradians.

- Laboratoire de Physique des Particules Fondamentales de l'École Polytechnique - Paris;
- Laboratoire de l'Accélérateur Linéaire - Orsay (Paris Sud);
- Centre de Recherches Nucléaires - Strasbourg.

Le flux d'hypérons signés par le compteur DISC présente un maximum en fonction de l'impulsion. Pour 2×10^{11} protons incidents sur la cible par cycle d'accélérateur, on obtient un maximum de $80 \Sigma^-$ à 19 GeV/c et un maximum de $1 \Xi^-$ à 16 GeV/c. Pour des impulsions plus faibles, le flux décroît par suite de la désintégration plus importante des hypérons.

Dans une deuxième expérience, on a mesuré les sections efficaces totales $\Sigma^- p$ et $\Sigma^- n$ à 19 GeV/c. Ces valeurs sont obtenues en mesurant l'atténuation d'un faisceau de Σ^- à travers une cible d'hydrogène et de deutérium liquides de un mètre de long. Un premier compteur DISC mesure le flux incident en amont de la cible, et un deuxième le flux transmis en aval de la cible. Des chambres à fils proportionnelles de haute résolution mesurent l'angle des hypérons avant et après la cible avec une précision de 0,7 milliradians. Pour contrôler la méthode, on a effectué des mesures des sections efficaces pp et pn à 19 GeV qui sont en parfait accord avec des résultats antérieurs.

De telles mesures permettent un test du modèle des Quarks. La prédiction de ce modèle est bien vérifiée pour la section efficace $\Sigma^- p$, moins bien pour la section efficace $\Sigma^- n$ (où la section efficace est trop faible de 2,5 écarts standards).

Désintégrations leptoniques d'hypérons

L'étude des désintégrations leptoniques des hypérons fait actuellement l'objet d'une expérience (fig. 2) qui devrait donner les nombres suivants d'événements :

$$\begin{aligned} 3000 \Sigma^- &\rightarrow n + e^- + \bar{\nu} \\ 100 \Xi^- &\rightarrow \Lambda + e^- + \bar{\nu} \\ 7 \Xi^- &\rightarrow \Lambda + e^- + \bar{\nu} \end{aligned}$$

L'accent est mis sur les deux premières désintégrations pour lesquelles, par

cette méthode, le nombre d'événements déjà observés dans l'ensemble des autres expériences est notablement augmenté. Ici encore un compteur DISC et deux chambres à fils placées de part et d'autre, alignent un Σ ou un Ξ et définissent sa direction avant désintégration. Cette dernière est observée dans une première chambre à streamers de 3 m de longueur. Un compteur Cerenkov à seuil de 2 m de long suit la première chambre à streamers. Ce compteur n'est sensible qu'aux particules ayant une vitesse supérieure au seuil et permet de ne retenir que les désintégrations qui produisent un électron. Une seconde chambre à streamers de 1,5 m de longueur est placée dans un gros aimant qui donne un champ magnétique de 0,8 Tesla. L'impulsion des particules chargées produites dans la désintégration est ainsi mesurée avec une précision de quelques centièmes.

Environ 5 mètres au-delà de l'aimant d'analyse, un détecteur à neutrons permet de localiser l'apex de la gerbe que le neutron y développe et de déduire ainsi l'angle du neutron émis dans la désintégration.

Les chambres à streamers et le détecteur à neutrons sont photographiés par deux caméras. Les films sont mesurés à Paris par un appareil automatique HPD.

Cette expérience permet de tester la validité des prédictions du modèle de Cabibbo pour les interactions faibles. Les constantes de couplages des courants axiaux et vectoriels se reflètent dans la densité des événements dans le diagramme de Dalitz. A partir de la distribution observée pour la désintégration $\Sigma^- \rightarrow n + e^- + \bar{\nu}$, on mesurera le module du rapport des constantes de couplage g/g_8 avec une précision accrue. Le signe de ce rapport n'est pas encore connu expérimentalement, mais on espère le déduire de la forme

du spectre en impulsion de l'électron. Dans la désintégration :

$$\Sigma^- \rightarrow \Lambda + e^- + \bar{\nu}$$

la mesure de la polarisation longitudinale du Λ permettra de confirmer que le courant correspondant est purement axial.

Trois domaines de recherches

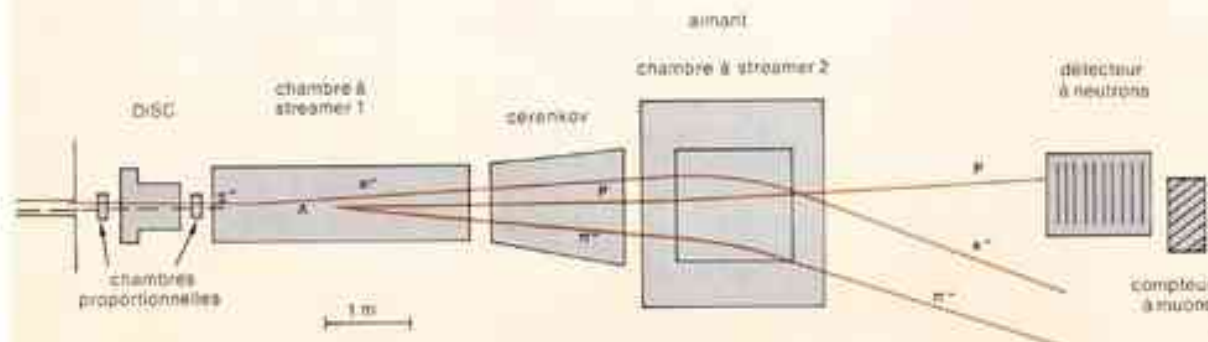
Le programme expérimental doit se poursuivre dans trois directions :

- Recherche des particules Ω^- . Celle-ci permettra de mesurer la section efficace de production. Seule une limite supérieure est actuellement connue. Jusqu'à maintenant le nombre total d'événements comportant une désintégration d' Ω^- enregistrés en chambre à bulles n'est que de 30 environ. Si le nombre d' Ω observés le permet on pourra aussi déterminer les rapports de branchement dans les trois canaux principaux ΛK^- , $\Xi^- \pi^0$, $\Xi^- \pi^-$.

- Recherche du mode de désintégration $\Xi^- \rightarrow n \pi^-$ interdit par la règle de sélection du nombre quantique d'étrangeté S ($\Delta S \leq 1$). En même temps, on augmentera le nombre des désintégrations $\Xi^- \rightarrow \Lambda + e^- + \bar{\nu}$ observées jusqu'ici.

- Programme d'étude des interactions $\Sigma^- p$, et plus particulièrement de la diffusion élastique et des résonances Y^* produites dans l'échange du poméron, ces résonances et leurs modes de désintégration étant encore très mal connus. Des expériences portant sur les désintégrations leptoniques d'hypérons et sur les interactions $\Sigma^- p$ sont aussi en cours à Brookhaven.

Un projet de faisceau d'hypérons à plus haute énergie est étudié pour l'accélérateur de 300 GeV du Supercern. L'accroissement d'énergie apporte un gain substantiel dans les flux de particules et permettra encore de nouvelles expériences.



Le montage expérimental employé pour détecter les désintégrations leptoniques comporte notamment deux chambres à streamers entre deux plaques distantes de 20 cm où applique l'effet de suite après le passage d'une particule ionisante une haute tension de l'ordre de 600 kV. A partir des centres d'ionisation créés dans le gaz (pression atmosphérique) par la particule une série d'éclaves se développe, formant des "dards" (ou "streamers") que l'on photographie, ce qui permet de localiser la trajectoire de la particule.

La séquence des opérations est la suivante : identification d'un Σ^- ou d'un Ξ^- par le DISC et localisation de leurs trajectoires par deux chambres proportionnelles situées de part et d'autre. Le compteur Cerenkov à seuil permet de n'observer que les désintégrations leptoniques (présence d'un électron). La deuxième chambre à streamer est placée dans un champ magnétique, ce qui permet la mesure précise des impulsions des particules chargées produites dans la désintégration. On a représenté à titre d'exemple la désintégration $\Sigma^- \rightarrow \Lambda + e^- + \bar{\nu}$. Le Λ étant neutre n'est pas directement observable mais, son parcours étant faible, il est observable dans la chambre I par ses produits de désintégration p et π^- . Dans le cas des désintégrations donnant un neutron, on peut déterminer la direction du neutron à l'aide d'un détecteur situé à 5 mètres derrière l'aimant d'analyse.

spectrométrie nucléaire à haute résolution avec les protons de 1 GeV du synchrotron saturne

Une expérience classique de physique nucléaire consiste à envoyer un projectile (proton ou autre noyau léger) de quelques MeV (ou quelques dizaines de MeV) sur un noyau et à mesurer le spectre en énergie et la distribution angulaire des projectiles diffusés. La physique nucléaire utilisant des projectiles de l'ordre du GeV est par contre un domaine nouveau qui n'a donné lieu jusqu'ici qu'à quelques explorations.

Que peut-on en attendre ? Rappelons que la longueur d'onde associée à une particule diminue lorsque l'énergie augmente. Il faut donc des particules de très grande énergie pour pouvoir observer les phénomènes liés à de courtes distances (de l'ordre des distances entre deux nucléons du noyau). Il faut de même des particules incidentes de grands moments si l'on veut connaître les distributions des moments des nucléons dans le noyau.

Spectrométrie à haute résolution

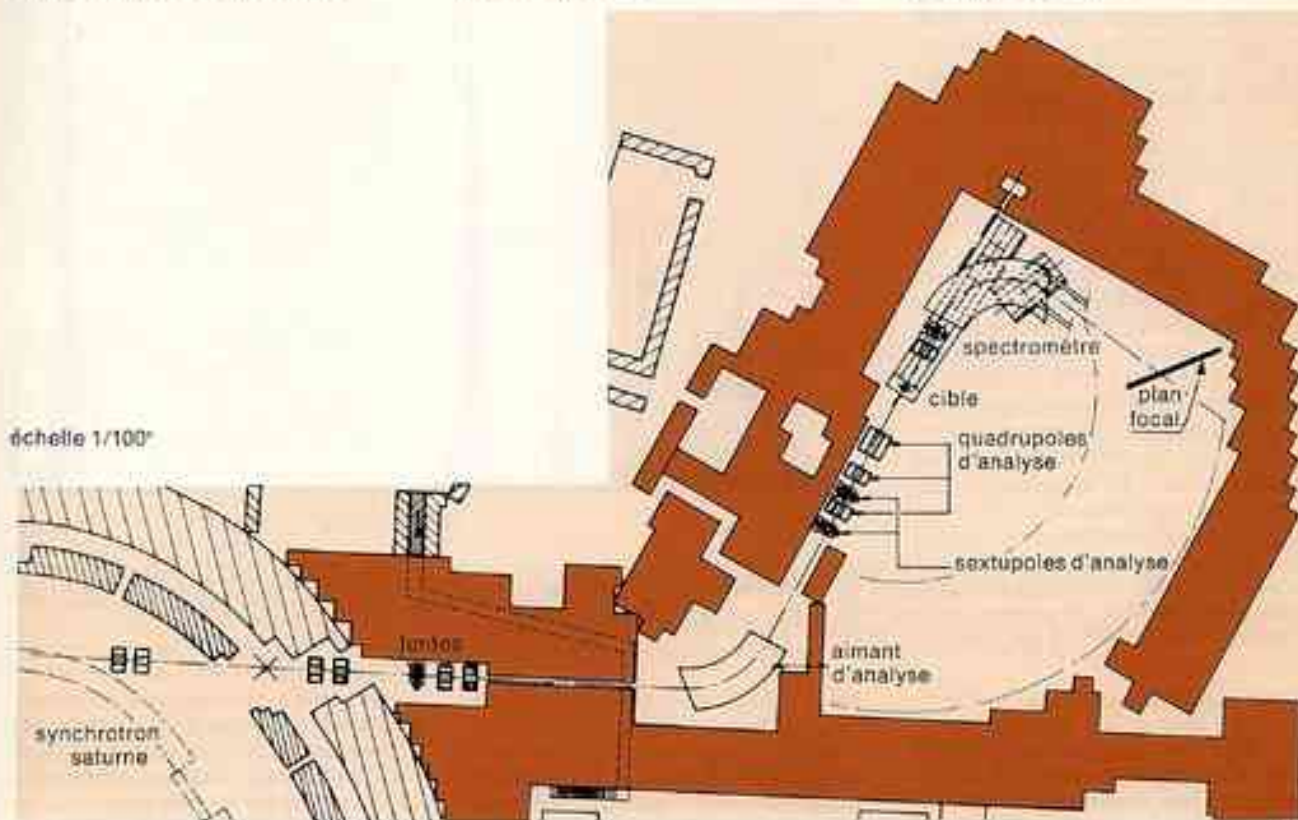
L'ensemble de spectroscopie nucléaire à haute résolution qui vient de donner ses premiers résultats a été conçu à Saclay pour pouvoir spécifier les états quantiques du noyau bien que l'énergie des projectiles soit de l'ordre de 1 GeV. Cela nécessite une résolution de l'ordre de 100 keV permettant de séparer un grand nombre d'états excités des noyaux. Cette résolution a été atteinte. Construit dans les années 1950 pour les besoins de la physique des particules, le synchrotron Saturne se révèle aujourd'hui très bien adapté à la physique nucléaire : son énergie est variable, la définition en énergie du faisceau est excellente.

- Service de Physique Nucléaire Moyenne Énergie (CIA - Saclay).
- Département du Synchrotron Saturne (CEA - Saclay).
- Centre de Recherches Nucléaires (CNRS - IN 2 P3) - Strasbourg.

(± 300 keV à 1 GeV) et l'étalement du temps de déversement des protons permet de ne pas saturer les détecteurs.

Rappelons le principe de cette expérience. Si un proton d'énergie E est envoyé sur un noyau cible et s'il excite un niveau E_1 , il ressortira avec une énergie légèrement différente $E' = E - E_1$. L'énergie E' est mesurée par un spectromètre magnétique et c'est pour séparer les niveaux E_1 qu'il faut avoir une excellente résolution en énergie.

Pour obtenir cette résolution il faut tout d'abord concevoir des méthodes qui permettent de compenser les effets parasites tels que la largeur en énergie du faisceau incident ou la variation en fonction de l'angle des énergies des particules sortant des cibles. Si ces effets parasites n'étaient pas compensés la résolution atteinte de 100 KeV n'aurait pu être meilleure que 2 ou 3 MeV.

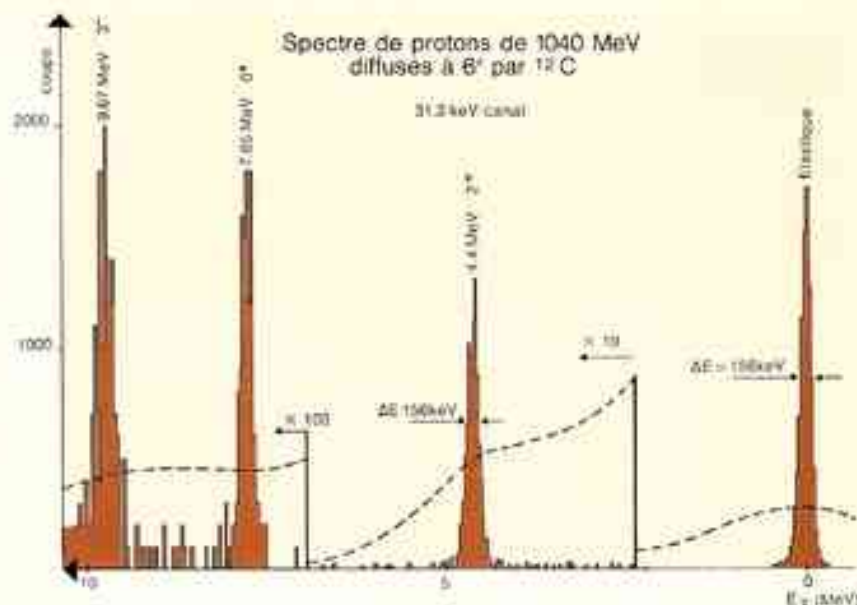


Le faisceau de protons extrait du synchrotron Saturne passe d'abord dans l'aimant d'analyse avant d'être focalisé sur la cible. Les particules diffusées par celle-ci passent dans le spectromètre et sont focalisées dans son plan focal. Le spectromètre tourne autour de la cible ce qui permet de faire varier l'angle de diffusion observé.

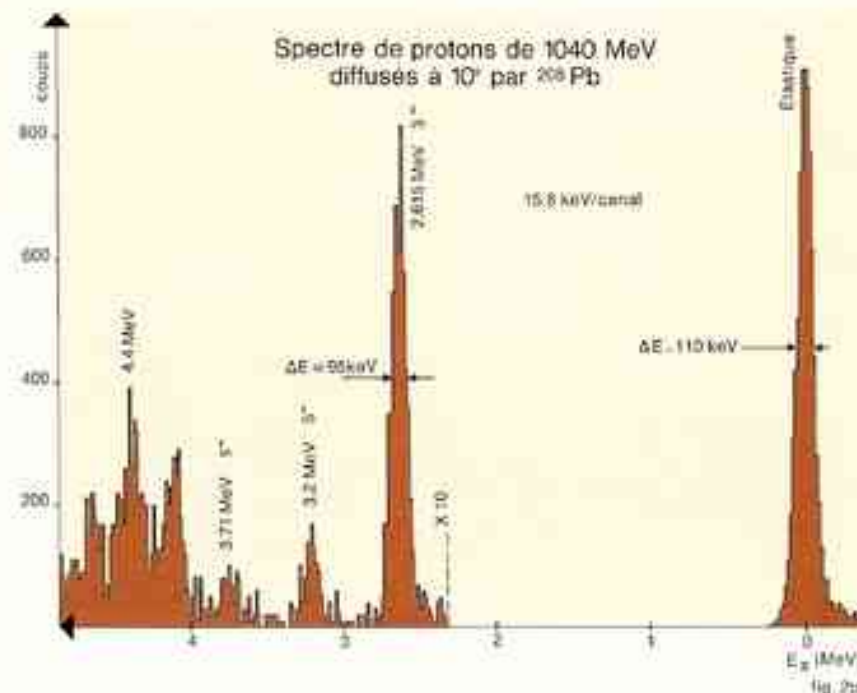
Montrons par exemple comment on corrige l'imprécision due aux variations d'énergie incidente. On utilise un ensemble de deux aimants dont l'un (analyseur) est situé sur le trajet des particules avant le choc sur la cible et l'autre (spectromètre) est situé sur le trajet des particules diffusées. Si les dispersions des aimants sont égales et de signe opposé, l'ensemble est achromatique. Les variations de l'énergie incidente n'auront pour résultat aucun déplacement dans le plan focal du spectromètre puisque l'analyseur et le spectromètre agissent sur elles. Par contre les variations d'énergie (E_x) apparaissant dans la cible par réaction nucléaire sont analysées par le seul spectromètre et provoquent un déplacement de trajectoire à sa sortie, ce qui permet de les identifier.

Mais ces méthodes de compensation exigent pour être réalisées que l'on définisse avec précision les caractéristiques des aimants. On se sert pour cela de programmes de calculs d'optique ionique qui permettent la détermination très précise des trajectoires des particules. Il reste à réaliser des aimants qui présentent les caractéristiques précises désirées sans pour autant exiger des tolérances d'usinage onéreuses. On obtient ce résultat en effectuant des mesures précises de champ magnétique sur l'aimant et en corrigeant par adjonction sur les culasses de pièces de fer minces, faciles à usiner ; il faut noter qu'il n'est pas nécessaire d'imposer une valeur déterminée du champ magnétique en tous points, seule l'intégrale le long des trajectoires étant essentielle. Le schéma général des aimants (fig. 1) montre une certaine complexité des divers éléments. Le spectromètre lui-même à un angle solide de $2 \cdot 10^{-5}$ stéradians. Il pèse environ 100 tonnes, ce qui est peu pour analyser des moments jusqu'à 2 GeV/c (des réalisations analogues pèsent 5 à 600 tonnes). Il tourne autour de la cible avec facilité sur des coussins d'air.

Grâce aux "compensations" mentionnées plus haut et aux aberrations négligeables de l'optique, la résolution ne dépend que de la largeur des fentes objets (0,5 mm environ ; objet virtuel) et de la précision avec laquelle on peut localiser, dans le plan focal, les particules qui correspondent à un niveau final du noyau cible. Cette localisation s'effectue



On porte ici le nombre de particules diffusées sur un noyau de ^{12}C , pour une position angulaire donnée, en fonction de l'énergie d'excitation E_x qu'elles ont laissée dans la cible. Pour $E_x = 0$, on obtient le pic de diffusion élastique visible à la droite de la figure. Pour $E_x \neq 0$, on obtient plusieurs pics d'excitation inélastique beaucoup plus petits (remarquer les changements d'échelle verticaux) correspondant aux divers niveaux excités de ^{12}C . Les niveaux d'énergie du ^{12}C sont très bien séparés. La courbe en pointillés montre que cela n'était pas le cas au cours des expériences antérieures où la résolution était de l'ordre de 2 à 3 MeV.



On porte ici le spectre en énergie d'excitation E_x pour un noyau lourd tel que le ^{208}Pb ; les niveaux γ sont mieux séparés que dans le cas du ^{12}C ce qui permet d'effectuer l'analyse angulaire dont il sera question plus loin (fig. 3).

tue à l'aide d'un détecteur de particules original, la chambre à migration. Le principe d'une telle chambre est le suivant : les électrons produits dans le gaz spécial par l'ionisation d'une particule se déplacent dans un champ électrique constant vers un compteur proportionnel situé à l'extrémité. Ils y arrivent après un temps qui dépend

linéairement de la position initiale des électrons. Le temps zéro est défini par un ensemble de scintillateurs. La mesure permet une localisation avec une précision de 0,5 à 1 mm environ selon les temps de migration ce qui est très satisfaisant. Finalement, l'intersection avec le plan focal est faite, en ligne, par ordinateur.

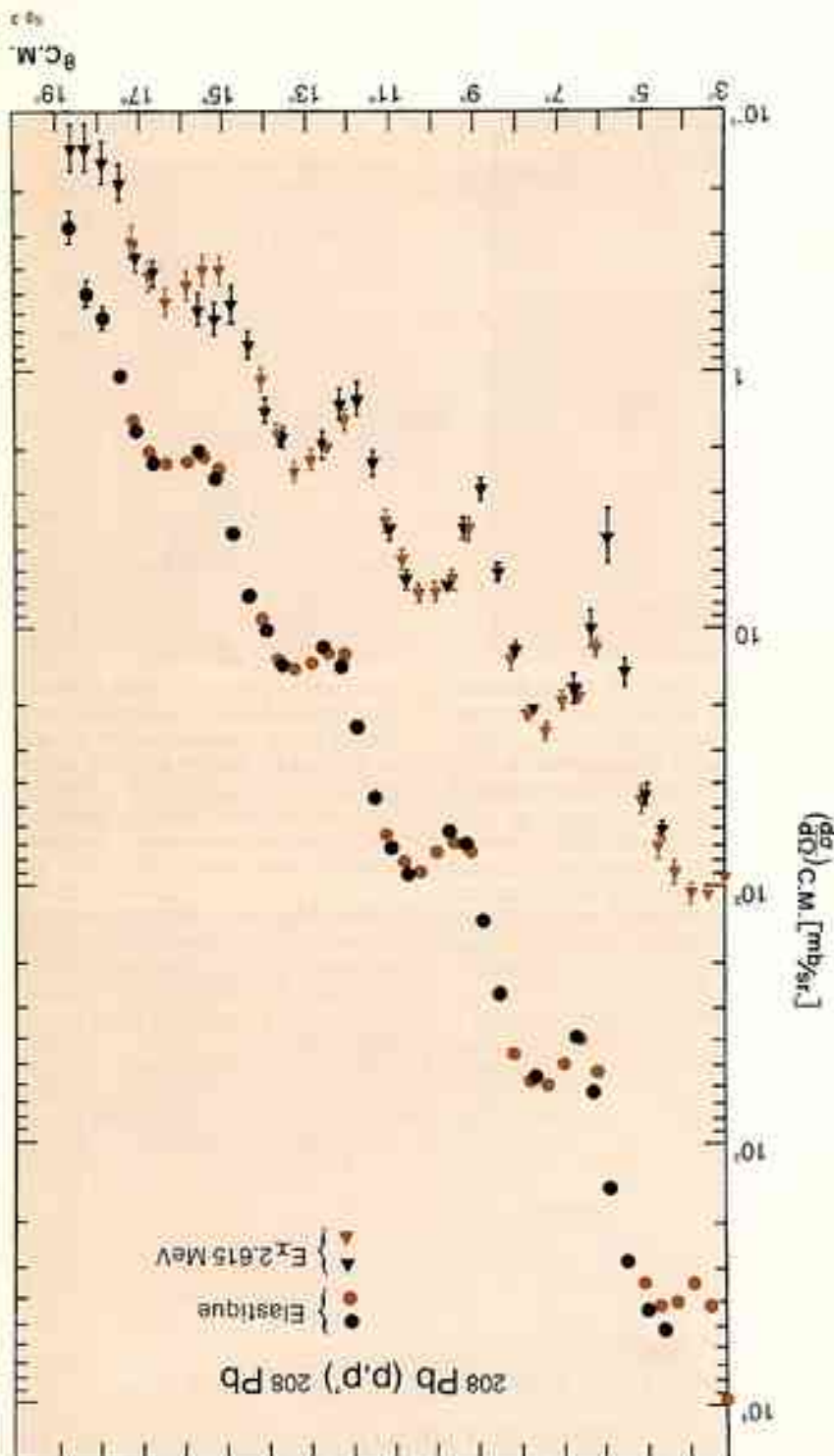
Les premiers résultats

La figure 2a présente, à titre d'exemple, un spectre d'énergie obtenu dans la diffusion de protons de 1 GeV sur une cible de ^{12}C .

On observe clairement les divers niveaux. Des expériences antérieures, à Brookhaven en 1967, donnaient une résolution de l'ordre de 2 MeV, avec des angles solides environ dix fois plus faibles. Des expériences très récentes de Gatchina (URSS) ont permis d'obtenir une résolution de l'ordre de 2 MeV et la séparation des niveaux y est donc souvent problématique ainsi que l'on peut s'en rendre compte par l'allure de la courbe en pointillés. La figure 2b montre que la résolution permet d'établir la séparation de niveaux excités d'un noyau lourd tel que ^{208}Pb .

Tous ces niveaux sont en fait bien connus depuis longtemps et la véritable nouveauté de ces expériences apparaît sur la figure 3 : on examine ici la distribution angulaire des protons émis en séparant la contribution du niveau fondamental (diffusion élastique) et du premier niveau excité du ^{208}Pb .

On remarque le caractère oscillatoire de la distribution angulaire dont l'interprétation théorique fait intervenir une théorie de la diffraction analogue à celle qui intervient en optique (le noyau est considéré comme un disque noir) ; c'est ici qu'interviennent les principes aux grands angles, donc aux grands moments transférés, que l'on attend à voir un comportement différent de celui apporté par la spectro-métrie nucléaire classique. De même, l'interaction du proton incident avec un seul nucléon du noyau devrait permettre de vérifier quantitativement une théorie de cette diffraction nucléaire qui la décrit comme une suite de diffractions multiples sur les nucléons individuels. Les écarts éventuels à cette théorie permettront peut-être de mieux comprendre les corrélations à courte distance entre les nucléons.



Un a cette fois représente la section efficace de diffusion en fonction de l'angle pour la diffusion élastique et celle correspondant au premier état excité J^π à 2,615 MeV du ^{208}Pb (les deux premiers pics de la figure 2b). On remarque le caractère oscillatoire des distributions angulaires dont l'interprétation fait intervenir une théorie de la diffraction (le noyau est considéré comme un disque noir) analogue à celle qui intervient en optique.

Les calculs sont assez complexes et sont actuellement en cours. L'essentiel est que ces données étaient hors de portée il y a seulement deux ans. Il s'agit donc d'un nouveau domaine expérimental plein de promesses. Cet enjeu est ainsi encore le champ des expériences possibles.

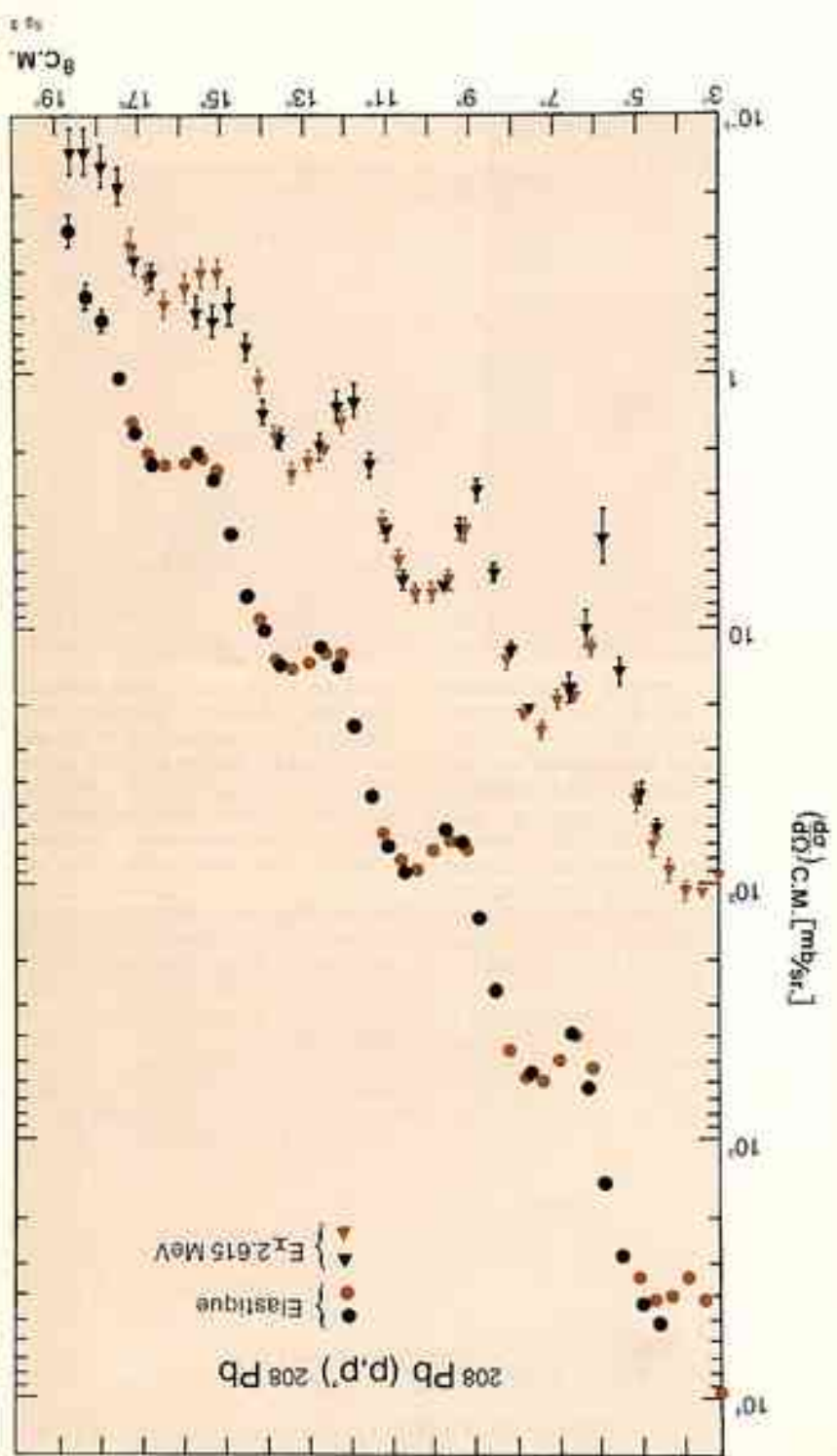
Les premiers résultats

La figure 2a présente, à titre d'exemple, un spectre d'énergie obtenu dans la diffusion de protons de 1 GeV sur une cible de ^{12}C .

On observe clairement les divers niveaux. Des expériences antérieures, à Brookhaven en 1967, donnant une résolution de l'ordre de 2 MeV, avec des angles solides environ dix fois plus faibles. Des expériences très récentes de Gatchina (URSS) ont permis d'obtenir une résolution de l'ordre de 2 MeV et la séparation des niveaux y est donc souvent problématique ainsi que l'on peut s'en rendre compte par l'allure de la courbe en pointillés. La figure 2b montre que la résolution permet également la séparation de niveaux excités d'un noyau lourd tel que ^{208}Pb .

Tous ces niveaux sont en fait bien connus depuis longtemps et la vérification de ces expériences apparaît sur la figure 3 : on examine ici la distribution angulaire des protons émis en séparant la contribution du niveau fondamental (diffusion élastique) et du premier niveau excité du ^{208}Pb .

On remarque le caractère oscillatoire de la distribution angulaire dont l'interprétation théorique fait intervenir une théorie de la diffraction analogue à celle qui intervient en optique (le noyau est considéré comme un disque noir). C'est ici qu'interviennent les principaux avantages des hautes énergies : c'est aux grands angles, donc aux grands moments transférés, que l'on s'attend à voir un comportement différent de celui apporté par la spectro-métrie nucléaire classique. De même, l'interaction du proton incident avec un seul nucléon du noyau devrait permettre de vérifier quantitativement une théorie de cette diffraction nucléaire qui la décrit comme une suite de diffractions multiples sur les nucléons individuels. Les écarts éventuels à cette théorie permettront peut-être de mieux comprendre les corrélations à courte distance entre les nucléons.



Un a cette fois représente la section efficace de diffusion en fonction de l'angle pour la diffusion élastique et celle correspondant au premier état excité $J^\pi = 2, 615 \text{ MeV}$ (les deux premiers pics de la figure 2b). On remarque le caractère oscillatoire des distributions angulaires dont l'interprétation fait intervenir une théorie de la diffraction (le noyau est considéré comme un disque noir) analogue à celle qui intervient en optique.

Les calculs sont assez complexes et sont actuellement en cours. L'essentiel est que ces données étaient hors de portée il y a seulement deux ans. Il s'agit donc d'un nouveau domaine expérimental plein de promesses. Cet enjeu est ainsi encore le champ des expériences possibles.

mise en évidence de noyaux "critiques" ou "supermous" loin de la zone des noyaux stables

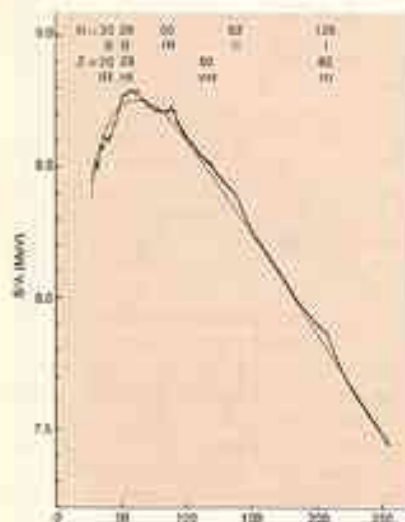
L'étude expérimentale des noyaux a montré depuis longtemps que certaines de leurs propriétés présentent des variations remarquables pour certaines valeurs du nombre Z de protons et N de neutrons. La figure 1 montre par exemple que les énergies de séparation des nucléons passent par des maxima comme cela se passe dans les atomes pour les potentiels d'ionisation des différents éléments.

Ces faits s'interprètent dans le cadre du célèbre "modèle des couches". Chaque nucléon se meut dans le noyau sous l'influence d'un potentiel moyen qui résulte des interactions combinées de tous les autres nucléons. Dans ce système quantique les orbites sont évidemment quantifiées d'où l'apparition,

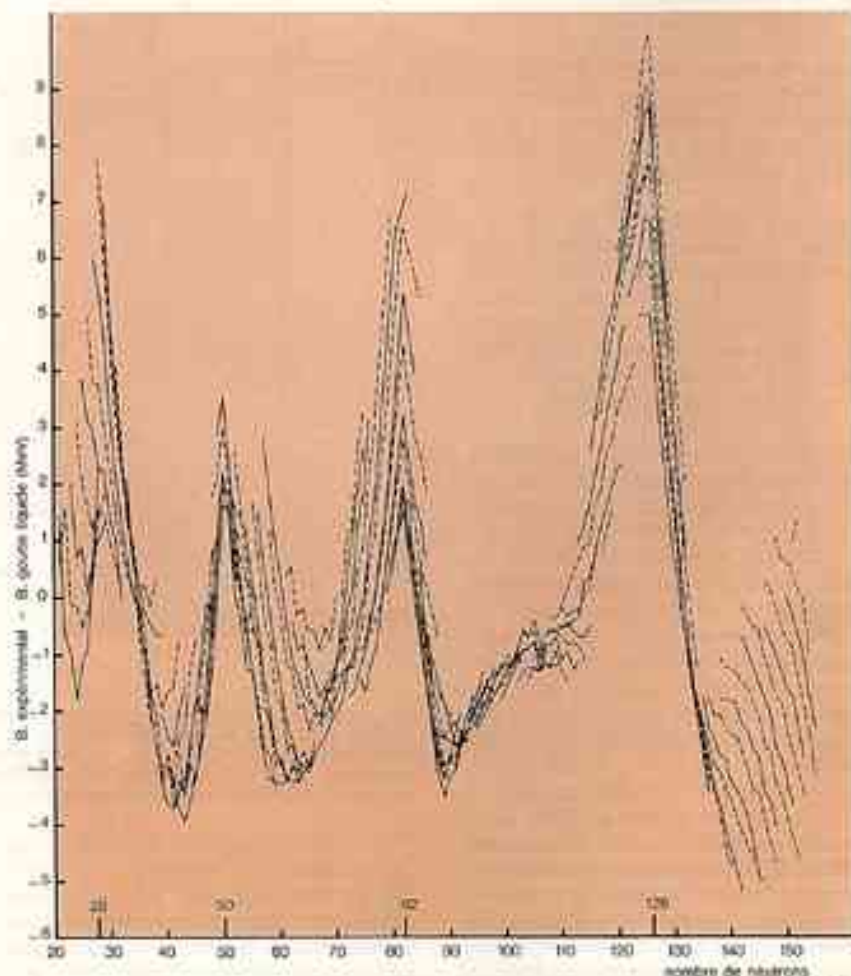
comme dans l'atome, de niveaux discrets, de couches fermées ("nombres magiques") etc... L'analogie avec l'atome est évidente mais des différences profondes existent. Par exemple, il y a dans l'atome un centre de force privilégié (le noyau qui contient toute la charge électrique positive) et ceci impose une symétrie sphérique. Rien de semblable dans le noyau où nul n'est privilégié et où le potentiel moyen peut prendre une forme différente de la sphère, si cette forme est avantageuse pour la collectivité en minimisant l'énergie totale de liaison.

Ceci se confirme de façon frappante dans la figure 2 où l'on porte un paramètre d'anisotropie en fonction du nombre N ou Z . Les noyaux "magi-

ques" sont effectivement sphériques mais, loin des couches fermées, les noyaux prennent la forme d'un sphéroïde allongé (cigare). Expérimentalement ces propriétés se manifestent par une répartition très anisotrope de charges électriques (moment quadrupolaire) et par l'allure des spectres d'énergie des niveaux excités. Un noyau allongé doit en effet présenter dans son spectre d'énergie des séquences de niveaux caractéristiques appelées bandes rotationnelles (par analogie avec les spectres des molécules). Comme l'espacement des niveaux dans une bande rotationnelle dépend en outre du "moment d'inertie" de ce corps en rotation, on en déduit également une information sur la rigidité des noyaux (fig. 4).



La courbe 1a représente l'énergie de liaison par nucléon des noyaux en fonction de leur nombre de masse. Le trait en couleur représente ce que donne pour cette grandeur un modèle semi-empirique qui assimile le noyau à une goutte liquide. On s'aperçoit que le comportement moyen des noyaux est remarquablement rendu par le modèle de la goutte liquide (à part les perturbations, au voisinage des nombres magiques, qui sont nettement visibles sur la courbe des valeurs expérimentales). L'effet sur l'énergie de liaison de la fermeture des couches est rendu encore plus net sur la figure 1b où l'on a porté la différence entre les valeurs expérimentales et les valeurs "goutte liquide" de l'énergie totale de liaison.



- Institut de Physique Nucléaire - Orsay (Paris-Sud)
- Collaboration Française au projet ISOLDE (CERN)

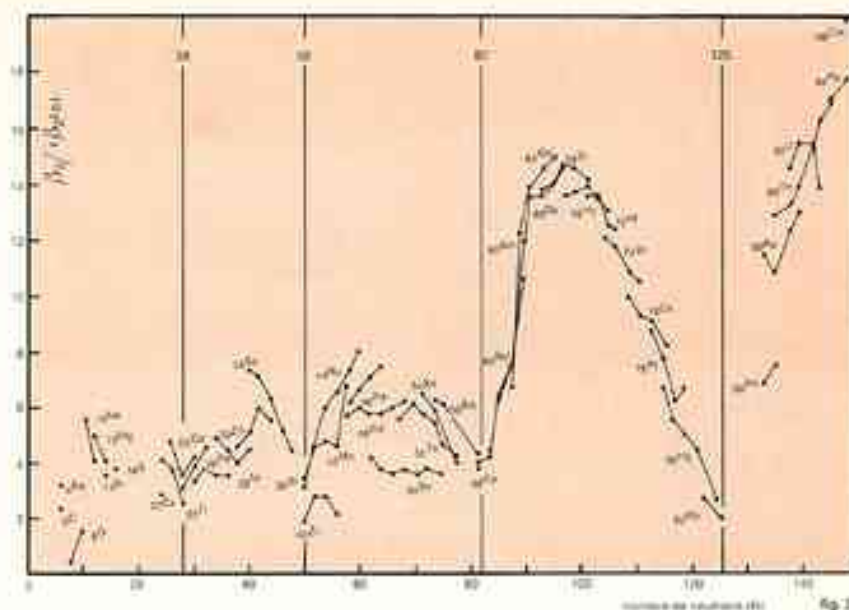
Si donc, par certains aspects, la structure du noyau ressemble à celle des atomes, d'autres propriétés au contraire rappellent celles de la matière condensée sous ses différentes phases. Les noyaux allongés à déformation permanente peuvent à la limite être considérés comme des bâtonnets solides.

Pour les noyaux "magiques", les nucléons sont quasi-indépendants, les seules contraintes étant l'organisation en couches et la symétrie sphérique : on peut s'imaginer une bulle de gaz comprimée de façon uniforme (au sein d'un liquide par exemple) dans lequel régnerait un système d'ondes stationnaires. Il existe des noyaux de forme intermédiaire entre la sphère et le bâtonnet, ceux qui présentent la forme d'un disque. Comme cette forme est instable, déformable, il est naturel de l'assimiler à une phase liquide. Naturellement, toutes ces "phases" sont quantiques et leurs propriétés "macroscopiques" sont liées aux états "microscopiques" des nucléons qui les composent. C'est ainsi que la tendance des neutrons et protons à se grouper par paires entraîne des propriétés voisines de la superfluidité.

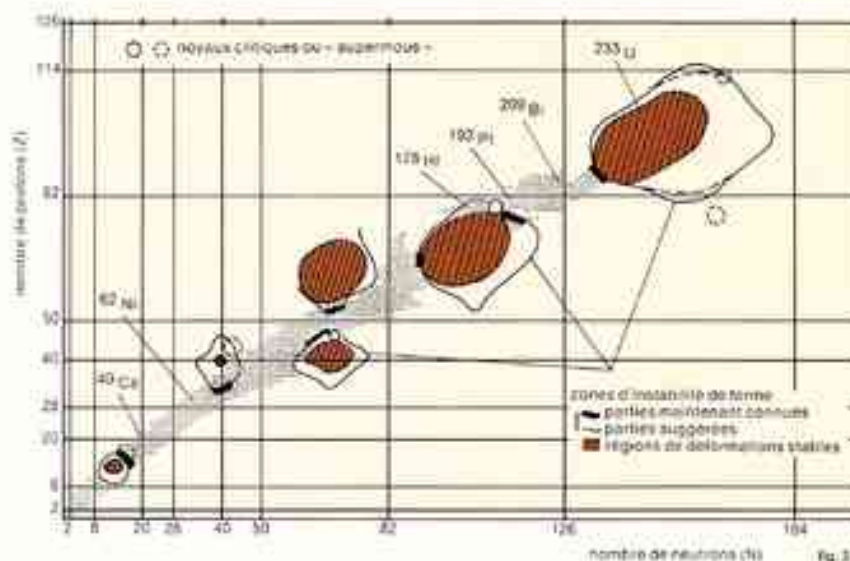
Comment mettre en évidence les transitions de phase dans les noyaux ?

C'est évidemment en étudiant en détail les caractéristiques des spectres des niveaux excités des noyaux que l'on obtient des renseignements sur leur structure. Pour observer les changements de phase éventuels, l'étude de noyaux isolés ne suffit pas. La démarche la plus féconde consiste à suivre une propriété (par exemple la position d'un niveau d'excitation) sur un grand nombre de noyaux différant par leur nombre N (ou Z) et à observer une discontinuité ou un extremum au moment du changement de phase. Cette démarche est d'ailleurs assez générale en physique nucléaire comme le montrent les exemples des figures 1 et 2.

Ce qui est assez inhabituel ici c'est la



En fonction du nombre de neutrons on porte un paramètre lié à la déformation. Les traits verticaux représentent les noyaux à couches fermées. On constate que c'est en milieu de couche, par exemple entre 82 et 126, que des noyaux prennent la forme d'ellipsoïdes allongés. Inversement les noyaux à couches "fermées" sont les moins déformés. Bien que cela apparaisse moins sur la figure, des zones de déformation aussi nettes existent dans les noyaux plus légers (voir fig. 3).



Les noyaux connus occupent une étroite zone du plan N-Z au voisinage de celle des noyaux stables. Il existe des zones au milieu de l'intervalle séparant deux couches fermées où les noyaux sont déformés de façon permanente (certaines d'entre elles, situées dans une zone accessible expérimentalement ont été bien étudiées). Entre les noyaux "solides" (à déformation permanente) et les noyaux sphériques "magiques" (analogues à une bulle de gaz comprimé) il existe des zones d'instabilité de forme où le noyau est très déformable (analogue à un liquide). Des points triples ou "noyaux-critiques" existeraient à la limite de ces trois "phases".

région du plan (N-Z) où il faut aller chercher ces noyaux. Nos connaissances actuelles reposent sur l'étude approfondie de noyaux situés dans une zone assez étroite du plan N-Z, au voisinage des noyaux stables existant dans la nature (figure 3). Or les calculs effectués depuis 5 ans aux Etats-Unis, en Suède, en URSS montraient que les régions où l'on pourrait observer éven-

tuellement tous ces changements de phase (et peut être même un point triple) se situaient aux confins de la zone connue. Plus précisément, il s'agissait de produire des noyaux lourds d'osmium, platine, mercure, plomb déficients en neutrons. Ces noyaux, instables par radioactivité β ont de courtes durées de vie ce qui était le principal obstacle à leur étude.



fig. 4

C'est en étudiant systématiquement la position de niveaux caractéristiques dans une série d'isotopes que l'on trouve des variations remarquables de propriétés. Dans le cas des platines représentés ici, l'écartement des niveaux de moment angulaire 0, 2, 4, 6 (en couleur) se rapproche des valeurs attendues pour une bande rotationnelle lorsque l'on se déplace vers des noyaux de plus en plus légers. C'est le cas par exemple pour $N \approx 104$, milieu de couche en neutrons où l'on sait qu'il s'agit d'une zone déformée (fig. 3 et 4). La position du premier niveau excité de moment angulaire nul (0^+) s'effondre entre les masses 188 et 186. L'inversion des positions relatives du premier niveau 4^+ et du deuxième niveau 2^+ (voir flèche) est caractéristique d'un point "critique" et s'interprète comme un changement de phase.

L'étude a été menée principalement par le groupe de spectroscopie nucléaire "en ligne" d'Orsay. On bombarde de l'or ou du plomb avec un faisceau de protons de 155 MeV (Orsay) ou de 600 MeV (CERN). On obtient ainsi par spallation des noyaux instables qui sont déficitaires en neutrons. Ils sont composés de nombreux isotopes qu'il faut séparer. On a utilisé pour cela un séparateur d'isotopes constitué d'une source ionisante, d'une accélération électrostatique, d'une déflexion magnétique et d'un collecteur.

Les rayonnements émis sont détectés à l'aide de cristaux de germanium et de silicium qui fournissent des impulsions électriques proportionnelles à l'énergie des rayonnements. Ces impulsions sont amplifiées, triées, codées et envoyées à un ordinateur qui les enregistre et les classe. Tout cela se fait "en ligne" au CERN. (Un tel ensemble de production et d'études "en ligne" est en cours de réalisation à Orsay). On détermine quels rayonnements sont émis en cascade et dans quel ordre, ainsi que leurs angles d'émission relatifs; de là on déduit les niveaux d'énergie et les moments cinétiques des états nucléaires (l'énergie des rayonnements γ est égale à la différence d'énergie entre deux états d'excitation nucléaire, comme c'est le cas pour les transitions électromagnétiques émises par un atome se désexcitant).

La discussion est ouverte sur la nature de ces noyaux

La densité des modes d'excitation près de l'état fondamental des noyaux croît brusquement dans la région des platines de masse 184-186 comme les niveaux indiqués dans la figure 4 en donnent un aperçu. Il y a en particulier dans les noyaux de platine de masse 184 et 186 un "effondrement" de l'énergie du premier niveau excité de moment angulaire nul (0^+) après inversion des positions relatives du premier niveau 4^+ et du deuxième niveau 2^+ . Ceci s'interprète comme le passage d'un état "liquide" à un état "mou" lorsqu'on se rapproche du milieu de couche en neutrons. Un effondrement des déplacements isotopiques (qui donnent la variation du rayon carré moyen des noyaux) a été trouvé dans les noyaux de mercure de la même région (masses 183 et 185) par une équipe allemande au CERN. Les mesures de radioactivité α et β effectuées par des chercheurs danois montrent également des "effondrements" pour des noyaux voisins. Ces diverses expériences effectuées à ORSAY et au CERN sont donc en bon accord et montrent bien qu'on se trouve à la limite de plusieurs phases dans les états de basse énergie des noyaux de $Z = 78-80$ et $N = 104-110$ et, plus particulièrement pour le platine de masse

186 où il y a peut-être coexistence de ces phases. Ce noyau semble bien être un des points triples recherchés.

Par contre sur l'interprétation (théorique (sur la nature exacte de ces "phases"), la discussion est très ouverte à l'heure actuelle. C'est ainsi que les chercheurs d'ORSAY ont avancé l'idée que ces noyaux sont extrêmement "mous" (analogues à des amibes); dans ces noyaux un grand nombre d'états quantiques ne seraient occupés qu'à temps partiel. Par contre les physiciens danois ont émis récemment une explication très hardie: on serait en présence de noyaux creux, à l'intérieur desquels existerait une cavité, ce qui reviendrait à dire que les états quantiques de moment angulaire nul ne seraient pas du tout occupés.

Mais il est possible qu'aucune de ces explications ne soit la bonne et que les interprétations soient encore modifiées par les résultats de nouvelles expériences en projet. Il est d'ailleurs probable qu'il existe d'autres régions du plan (N-Z) (voir figure 3) où de semblables phénomènes se manifestent. Ces aspects incontestablement nouveaux de la structure nucléaire pourraient jeter une lumière nouvelle sur un problème plus général en physique, celui des rapports entre les états "individuels" et "collectifs" dans les systèmes à N corps fortement couplés.

atomes
molécules
matière condensée



La lumière émise par un jet d'ions dans une expérience de "Beam-File"

études collisionnelles d'ions moléculaires

mais ceux-ci restent limités à un petit nombre de transitions.

Pour tourner cette difficulté, on utilise des méthodes dites "collisionnelles" qui donnent lieu à toutes sortes de spectroscopies. Leur dénominateur commun est la mesure de la distribution d'énergies des particules émergeant de collisions entre un faisceau de particules incidentes et une cible, qui peut d'ailleurs être constituée par un second faisceau. On entend ici par "particules" aussi bien des atomes des molécules, des ions, des électrons, des photons, etc. Nous allons décrire deux exemples d'application de ces méthodes collisionnelles :

- Photodissociation d'ions H_2^+ .
- Étude par coïncidences des excitations dans les collisions $He^+ - He$.

photodissociation d'ions H_2^+

L'étude des interactions entre la lumière et la matière demeure une branche très vivante de la physique. Or on sait peu de choses encore sur les spectres optiques des ions moléculaires à l'état gazeux, qu'ils soient sous forme de molécules ionisées ou de radicaux libres ionisés. Pourtant, certaines de ces espèces revêtent une grande importance en Astrophysique.

Pourquoi cette lacune ? En laboratoire, on ne peut produire des espèces ionisées pures que sous la forme de faisceaux que l'on transporte bien, certes, mais où la densité des particules est infime. De ce fait, les spectres d'absorption des ions moléculaires gazeux sont encore à peu près inaccessibles à l'expérience ; on peut enregistrer des spectres d'émission,

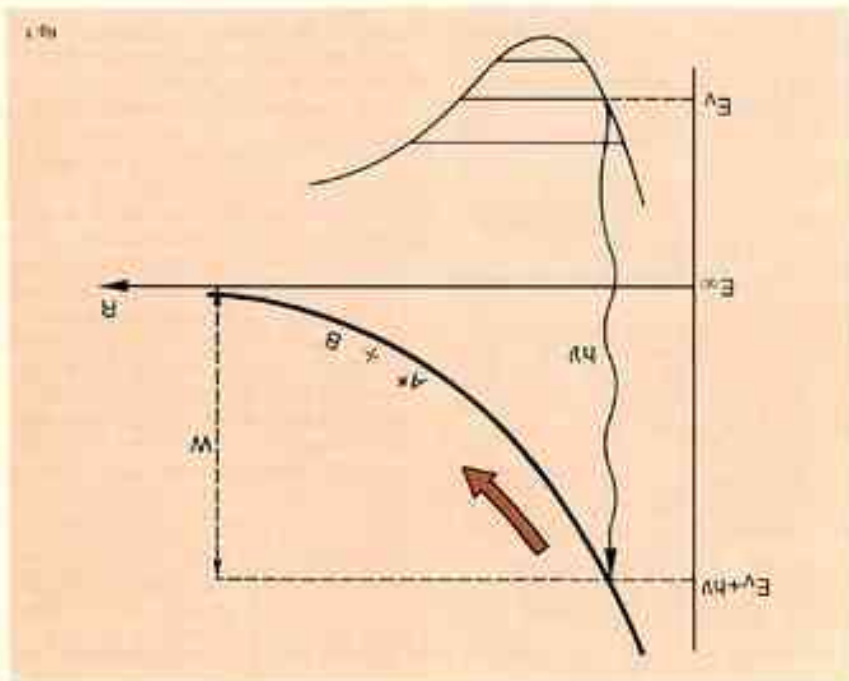
Une des dernières nées parmi les spectroscopies collisionnelles est la "spectro-photodissociation d'ions rapides", C'est l'analyse des vitesses des fragments qui se séparent lorsqu'un ion, composé de plusieurs atomes, est dissocié sous l'action d'un rayonnement lumineux.

Imaginer par exemple un faisceau d'ions positifs diatomiques AB^+ accélérés jusqu'à une énergie cinétique de plusieurs milliers d'électrons volts. Selon la façon dont ils ont été produits, ces ions moléculaires se trouvent à des niveaux d'énergie interne différents. En particulier, leur énergie E_i de vibration peut prendre diverses valeurs que l'on caractérise par l'nombre quantique $v = 0, 1, 2, \dots$ (fig. 1).

Mesure de l'énergie interne

Que se passe-t-il maintenant si on croise à angle droit ce faisceau d'ion avec un faisceau lumineux intense, monochromatique et parfaitement polarisé (le faisceau émis par un faisceau laser) ?

L'énergie interne de chaque ion AB^+ susceptible d'absorber un photon s'accroît de l'énergie $h\nu$ du photon (h est la constante de Planck, ν la fréquence de l'onde lumineuse). Si cette nouvelle valeur $E_f = E_i + h\nu$ de l'énergie interne de l'ion est supérieure à l'énergie qu'aurait les fragments A^+ et B^+ à distance infinie (E_∞), l'ion AB^+ "descend" spontanément vers ses produits de dissociation A^+ et B^+ .



variation, en fonction de la distance R , de l'énergie d'interaction $E(R)$ des particules A^+ et B^+ . La courbe du bas représente $E(R)$ lorsque l'atmosphère électronique du système A^+B^+ est dans son état fondamental. Les trois courbes schématisent les niveaux de vibration $v = 0, 1$ et 2 du H_2^+ . Lorsque cet ion absorbe un photon d'énergie $h\nu$, son atmosphère électronique est excitée et peut passer dans un état où $E(R)$ ne présente plus de minimum (coulée du haut de la figure qui correspond à une énergie repulsive). Juste après absorption du photon, on obtient un état d'énergie E_f qui est instable. Les deux particules A^+ et B^+ tendent alors à se repousser (R tend à augmenter) ; on aboutit donc à une dissociation de l'ion A^+B^+ . En deux particules A^+ et B^+ , ces dernières emportent une énergie cinétique totale W qui, comme le montre la figure, vaut $W = E_f - E_\infty$.

Ces derniers se séparent avec une énergie cinétique de translation relative W qui est égale à l'excès d'énergie disponible $E_f - E_\infty$ (fig. 1).

• Sous réserve que l'état électronique atteint par l'ion après absorption du photon soit "connecté" à $A^+ + B^+$ et non à $B^+ + A^+$.

• Laboratoire de Physique-Chimie des Rayonnements (Orsay) - Paris Sud.

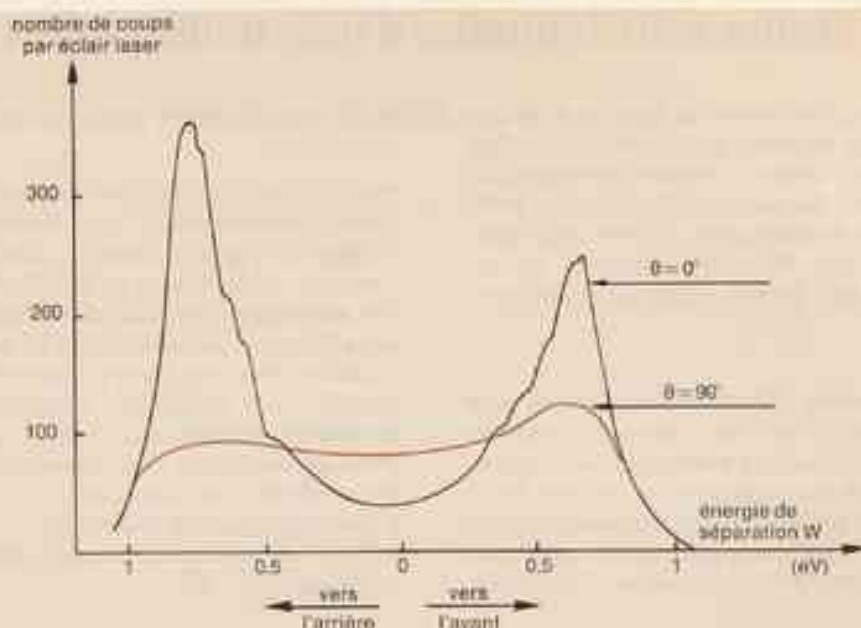


Fig. 2

Spectres d'énergie de séparation W des fragments H^+ et H de la photodissociation d'ions H_2^+ par le faisceau d'un laser à rubis. Les deux moitiés du spectre correspondent au départ du fragment H^+ soit vers l'avant, soit vers l'arrière, par rapport à la trajectoire des ions incidents. θ est l'angle que fait avec cette trajectoire la direction de polarisation de la lumière du faisceau laser.

Il faut évidemment choisir l'espèce AB^+ de telle sorte qu'elle puisse absorber le photon d'énergie $h\nu$; ceci exige qu'il existe à l'énergie E_i un état électronique, différent de l'état électronique fondamental, et accessible à partir de ce dernier compte tenu "des règles de sélection" optiques.

En supposant ces conditions satisfaites, comment le phénomène va-t-il pouvoir être observé ? S'il est difficile de déterminer directement l'énergie interne initiale E_i de l'ion AB^+ , par contre, il est facile de mesurer l'énergie cinétique T du fragment chargé A^+ à l'aide d'un analyseur électrostatique ou électromagnétique. Ainsi, de T on déduit W , de W on remonte à $E_i = W + E_{int}$ (E_{int} étant connu par la spectroscopie atomique) ; de E_i enfin, on redescend à $E_f = E_i - h\nu$.

Mais pourquoi avoir initialement accéléré les ions "parents" ? Parce que, grâce à la loi de composition des vitesses, si on observe les fragments A^+ se séparant dans la direction du faisceau d'ions parents (vers l'avant ou vers l'arrière), on peut mesurer W avec une précision d'autant plus grande que l'énergie cinétique initiale des ions parents est plus élevée. Cette précision, dans les cas favorables, peut théoriquement atteindre le dix millième d'électron-volt.

Ainsi, les spectres d'énergie de translation des fragments A^+ obtenus reflètent par leur structure les niveaux d'énergie E_i et les populations des états vibrationnels des ions AB^+ parents.

Sélection des ions moléculaires

Mais cette "photographie" du faisceau incident va même plus loin : lorsque le faisceau lumineux est polarisé, et c'est le cas pour un faisceau laser, les "règles de sélection" qui gouvernent l'absorption d'un photon par les ions AB^+ dépendent de l'orientation de l'axe moléculaire par rapport à la direction de polarisation de la lumière. Par exemple, les ions H_2^+ , issus de molécules d'hydrogène, absorbent d'autant mieux un photon que leur axe moléculaire a une orientation plus proche de la direction du vecteur lumineux. Ceci permet donc en principe d'effectuer une sélection des ions moléculaires selon leur orientation dans l'espace.

Par ailleurs, la dissociation de l'ion qui suit l'absorption du photon s'effectue en général beaucoup plus vite que sa rotation. La direction de l'axe moléculaire se retrouve donc comme direction de séparation des fragments. Si, comme cela a été dit plus haut, on "observe" les fragments émis dans la direction de la trajectoire du faisceau d'ions parents, le spectre des valeurs de l'énergie de translation W (donc de E_i) ne sera visible que si le faisceau lumineux est polarisé parallèlement au faisceau d'ions.

La figure 2 montre des spectres d'énergie de translation de protons H^+ qui résultent de la photodissociation d'ions H_2^+ rencontrant le faisceau d'un laser à rubis. La courbe obtenue pour une pola-

risation du faisceau laser parallèle au faisceau d'ion (courbe marquée $\theta = 0^\circ$) présente des échelons sur les deux pics correspondant respectivement à l'émission des ions H^+ vers l'avant et vers l'arrière. Ceux-ci reflètent la structure en niveaux vibrationnels E_i successifs des ions H_2^+ incidents. Ces niveaux ne sont pas encore bien résolus. Il s'agit de la première expérience de ce type où l'appareil a été réglé dans des conditions de basse résolution permettant d'obtenir l'intensité maximale du signal.

La courbe marquée $\theta = 90^\circ$, obtenue pour une polarisation du faisceau laser perpendiculaire au faisceau d'ions, est très différente de la précédente. Toute structure a disparu et le spectre s'est aplati. Théoriquement, il devrait se redresser si l'on produisait les ions parents H_2^+ avec une énergie de rotation telle qu'ils aient le temps de tourner avant de se dissocier, "oubliant" ainsi l'orientation qu'ils devaient avoir lors de l'absorption des photons.

Cette expérience ouvre la voie à une véritable spectroscopie des ions moléculaires, permettant de déterminer à la fois les niveaux vibrationnels de l'état électronique fondamental (ou d'un état métastable) et la forme des courbes de potentiel d'états dissociatifs. Cette technique peut en principe être utilisée pour n'importe quels ions moléculaires pourvu que l'on dispose de lasers opérant en régime pulsé, dans un domaine de longueurs d'onde plus courtes (c'est-à-dire émettant des photons d'énergie plus élevée) que le laser à rubis.

étude par coïncidence ion-photon des excitations dans les collisions $\text{He}^+ - \text{He}$

Comme nous l'avons vu plus haut, les méthodes classiques de spectroscopie d'absorption ou d'émission ne s'appliquent pas à l'étude d'ions atomiques AB^+ se trouvant dans des états électroniques peu stables ou même franchement dissociatifs (états où les deux atomes se repoussent quelle que soit la distance R qui les sépare). En effet, de telles espèces dont la durée de vie est très faible ne peuvent être obtenues qu'avec une densité infime. Par contre il est possible, partant des partenaires A^+ et B placés dans des états bien choisis, éventuellement excités, de les faire entrer en collision pour former transitoirement l'espèce AB^+ , cette "quasi molécule" se dissociant ensuite pour redonner A^+ et B (la succession de ces événements constitue la collision élastique $\text{A}^+ + \text{B}$). En réalité, il est évidemment beaucoup plus simple de partir de A^+ et B tous deux dans leur niveau fondamental et d'induire la formation de la molécule excitée $(\text{AB}^+)^*$ par la collision elle-même.

Dans quelles conditions et de quelle façon cette excitation par collision se produit-elle ? Nous allons voir que cela dépend essentiellement des propriétés de la molécule ionisée AB^+ .

Excitation par collision ion-atome à basse énergie

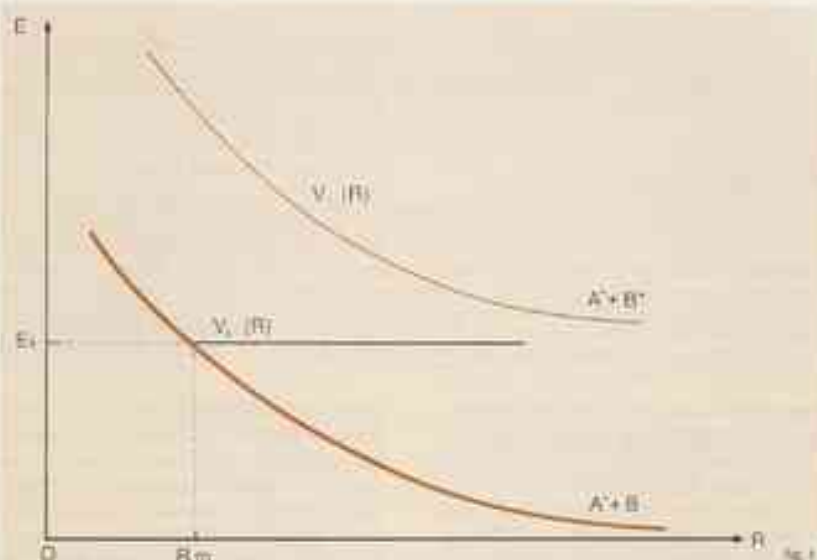
Considérons un ion A^+ et un atome B et supposons d'abord que leurs noyaux sont fixes, à une distance R . Les électrons du système ont alors un certain nombre de niveaux d'énergie, qui sont fonctions du paramètre R ; on désignera par $V_0(R)$ l'énergie du niveau fondamental, $V_1(R)$ celle du premier niveau excité, etc. (fig. 1). Lorsque R est très grand, ces niveaux correspondent à un ion A^+ et un atome B isolés, qui sont tous deux dans leur état fondamental si l'énergie est V_0 ; si l'énergie est V_1 , V_2 , etc., l'un est dans son état fondamental et l'autre dans son niveau excité. Pour les valeurs plus faibles de R , on n'a plus un ion A^+ et un atome B distincts, car les orbitales électroniques se recouvrent fortement, mais on est en présence d'un système global qui est la molécule ionisée AB^+ . Les états pour lesquels R étant fixé, les électrons ont une des énergies $V_0(R)$, $V_1(R)$, etc., sont appelés "états adiabatiques".

Que se passe-t-il maintenant dans une collision, c'est-à-dire lorsque R n'est pas constant ? Dans une collision, R

— Institut d'Électronique Fondamentale —
Orsay (Paris Sud)

part d'une valeur très grande qui diminue ensuite, ce qui conduit transitoirement à la formation d'une molécule AB^+ , puis augmente à nouveau après la collision, lorsque les deux partenaires s'éloignent. Les énergies $V_0(R)$, $V_1(R)$... jouent, pour les noyaux, le rôle d'une énergie potentielle à partir de laquelle on peut calculer leur mouvement. Appelons E l'énergie cinétique initiale des noyaux, lorsqu'ils sont à une distance R très grande. Si l'on avait une seule courbe de potentiel $V_0(R)$, la collision se déroulerait de la manière suivante : on suivrait par continuité l'état adiabatique obtenu pour chaque valeur de R , la valeur minimale R_m de R serait obtenue lorsque $E = V_0(R)$, puis R prendrait ensuite des valeurs de plus en plus grandes. L'état adiabatique obtenu après la collision correspondrait nécessairement à la courbe $V_0(R)$ et, par suite, on ne pourrait avoir que A^+ et B dans leur niveau fondamental (collision élastique).

En fait, il existe de nombreuses courbes de potentiel. Pour obtenir une excitation de l'atome ou de l'ion dans la collision, il faut, partant initialement sur la "voie" d'entrée associée à la courbe $V_0(R)$, "sauter" sur une autre courbe, par exemple $V_1(R)$; de cette façon, le système sort de la collision par une "voie" où l'atome B est excité. En d'autres termes, il faut avoir un "couplage" entre les différents états adiabatiques de la molécule AB^+ . Dans le cas où $V_0(R)$ et $V_1(R)$ sont des énergies très différentes pour toutes les valeurs de R (fig. 1), de tels couplages sont tout à fait négligeables lorsque les particules ont des énergies moyennes ($E = 100 \text{ eV}$). Cependant, dans certaines circonstances, ce couplage peut parfaitement se manifester. Un cas typique est celui où le niveau fondamental de AB^+ et le niveau excité ont même symétrie (Σ par exemple) et présentent un "croisement évité" à une certaine distance internucléaire R_c (fig. 2) : au voisinage de $R = R_c$, les énergies potentielles adiabatiques V_0 et V_1 sont très voisines et les courbes qui les représentent s'infléchissent sensiblement en ce point (elles "évitent" de se croiser). Dans une collision, R varie rapidement et ceci entraîne un couplage entre les états adiabatiques V_0 et V_1 lorsque $R = R_c$, de sorte que le système a une certaine probabilité de sauter d'une voie à l'autre. En conséquence,



Variations en fonction de la distance internucléaire R des énergies électroniques possibles pour le système $\text{A}^+ + \text{B}$, les noyaux étant supposés fixes. Pour les grandes valeurs de R , le niveau fondamental d'énergie $V_0(R)$ correspond à un ion A^+ et à un atome B isolés et tous deux dans leur niveau fondamental. Il en est de même pour le premier niveau excité d'énergie $V_1(R)$ à partir duquel l'un ou l'autre est excité (sur la figure, on a supposé que c'est l'atome, noté alors B^*). R_m représente la valeur minimale prise par R lors d'une collision entre A^+ et B avec une énergie cinétique initiale E_i .

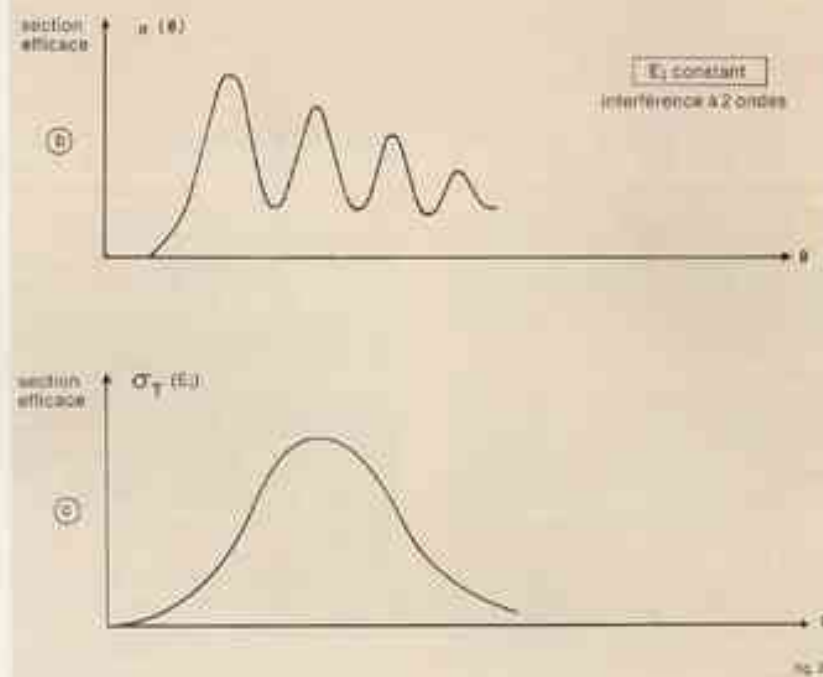
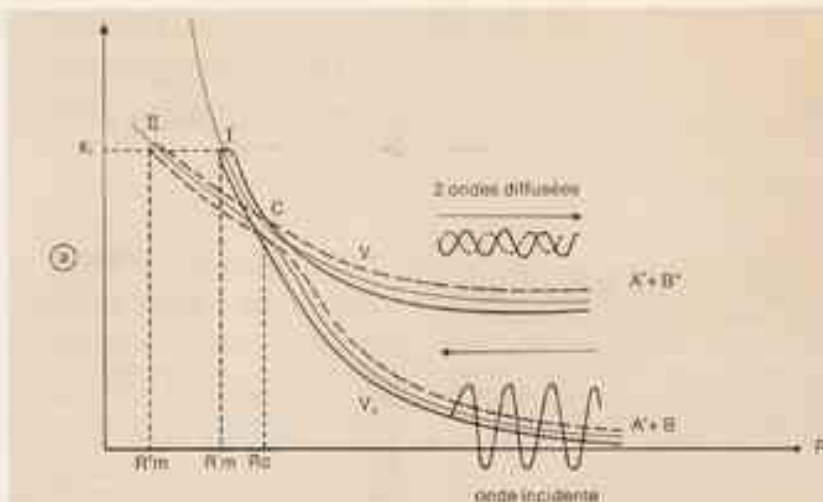
lorsque R passe par des valeurs voisines de R_c (ce qui est le cas si l'énergie initiale E_i est supérieure à $V_0(R_c)$), la collision inélastique

$A^+ + B \rightarrow A^+ + B^*$
est possible.

Effets quantiques d'interférence

Dans le mécanisme d'excitation qui a été envisagé, plusieurs "chemins" sont possibles pour le système. Ce dernier, partant initialement sur la voie d'entrée $A^+ + B$ avec l'énergie E_i , peut, lorsque $R = R_c$, sauter "à l'aller", sur l'autre courbe de potentiel : R diminue ensuite jusqu'à la valeur R_0 telle que $V_0(R_0) = E_i$, puis augmente jusqu'à l'infini (le système sort sur la "voie" donnant $A^+ + B$, chemin I de la fig. 2). Mais le système peut aussi sauter "au retour" d'une voie à l'autre : dans ce cas, R commence par diminuer jusqu'à R_0 tel que $V_0(R_0) = E_i$ et ce n'est que quand R repasse en augmentant par la valeur R_c que le système "saute" sur la voie donnant $A^+ + B^*$ (chemin noté II sur la fig. 2).

Classiquement, il faudrait ajouter les probabilités pour que le système passe soit par le chemin I, soit par le chemin II. En mécanique quantique, il en va autrement car le système peut suivre "à la fois" les deux histoires. En effet, à une particule est associée une onde et cette onde peut être diffusée en passant par l'un des chemins I ou II avec des amplitudes de probabilité complexes f_I et f_{II} : le chemin I donne dans la voie $A^+ + B$ une onde diffusée d'intensité $|f_I|^2$, le chemin II une onde d'intensité $|f_{II}|^2$, et si l'un de ces chemins était le seul, la probabilité d'excitation serait proportionnelle soit à $|f_I|^2$, soit à $|f_{II}|^2$. Mais l'onde dans la voie $A^+ + B^*$ est la somme des deux ondes cohérentes associées aux chemins I et II (fig. 2a) : la phase relative des deux ondes étant parfaitement déterminée, on ne peut pas simplement ajouter leurs intensités ; celle de l'onde totale sera donc $|f_I + f_{II}|^2$, et non $|f_I|^2 + |f_{II}|^2$. Comme dans une expérience d'interférences lumineuses, on aura donc des effets d'interférence entre les deux ondes, où leur déphasage, qui est lié à la différence du temps que met le système à parcourir les chemins I et II, joue un rôle crucial.



Excitation due à un croisement simple

Pour $R = R_c$, les deux courbes de potentiel en couleur V_0 et V_1 présentent un croisement évité (fig. 2a). Initialement, le système est constitué d'un ion A^+ et d'un atome B dans son niveau fondamental ; lorsque $R = R_c$, il peut "sauter" d'une courbe de potentiel à l'autre et, à la fin de la collision, l'atome B peut être excité. Deux chemins sont possibles pour le système, le chemin I représenté en trait plein et le chemin II représenté en trait pointillé. Chacun de ces chemins donne une onde diffusée. Des effets d'interférence sont possibles entre les deux ondes. Par suite (fig. 2b), la section efficace différentielle $\sigma(\theta)$ présente des oscillations en fonction de l'angle de diffusion θ . Cependant, ces oscillations disparaissent (fig. 2c) lorsqu'on s'intéresse à la section efficace totale σ_T .

Comment se manifestent, en pratique, de tels effets d'interférence ? On peut montrer que le déphasage entre les deux ondes dépend, à énergie E_i donnée, de la direction dans laquelle sont diffusées les particules après collision. Par suite, si l'on mesure la section efficace d'excitation en fonction de l'angle θ de diffusion (section efficace différentielle), on observe des oscillations dues au fait que le déphasage varie (fig. 2b). De même si on mesure, pour un angle de diffusion donné, les varia-

tions de la section efficace en fonction de l'énergie E_i , on a des oscillations sur la courbe obtenue. Toutefois, si c'est la section efficace totale à laquelle on s'intéresse, les oscillations disparaissent : on comprend bien que les variations rapides de la courbe 2b s'annulent lorsqu'on fait la somme des sections efficaces pour toutes les valeurs de θ . Dans une mesure de la section efficace totale en fonction de l'énergie E_i , on obtient donc une courbe qui a l'allure de celle de la fig. 2c.

Mécanisme d'excitation de Rosenthal

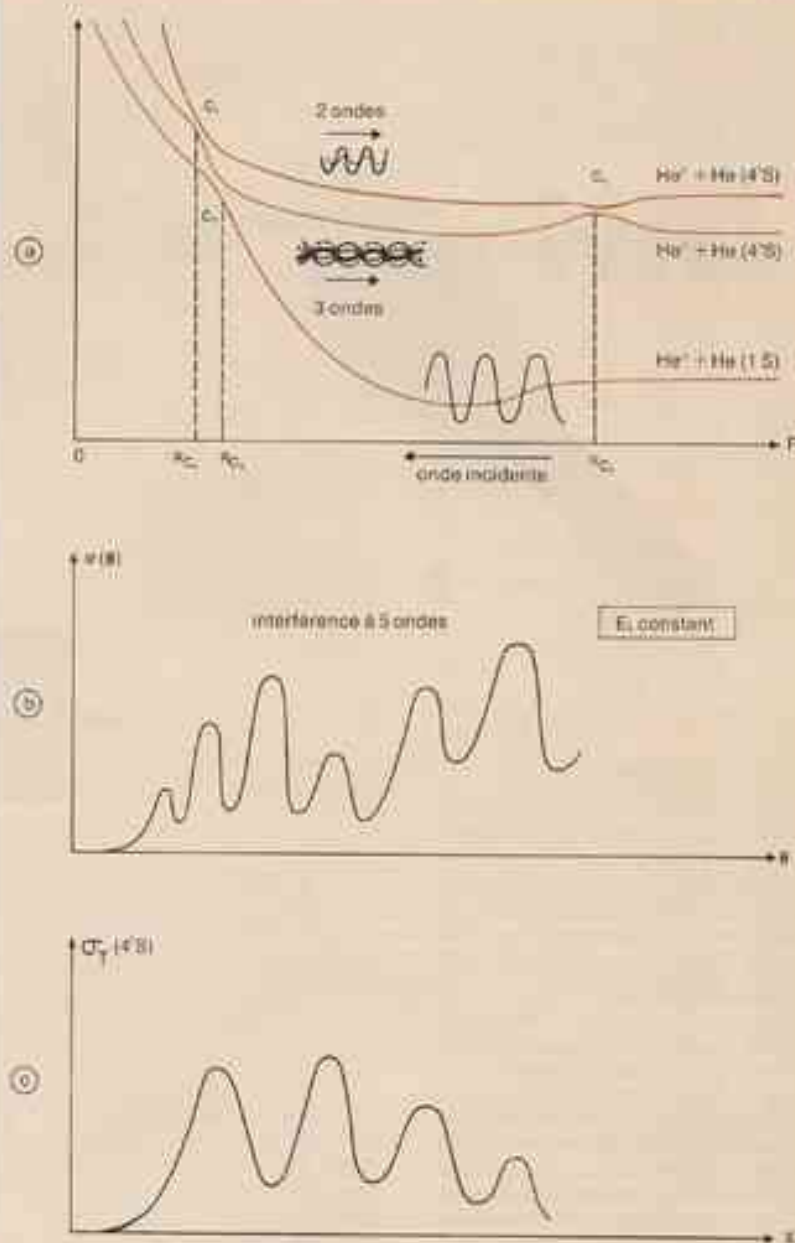
Il existe des cas où l'on ne peut pas se limiter à considérer deux courbes de potentiel, c'est-à-dire deux chemins possibles pour le système. Prenons par exemple le cas où la collision étudiée est



Le niveau excité 4S de l'hélium se compose de deux sous-niveaux multiplets, 4'S et 4'S, dont la différence d'énergie est relativement faible. Rosenthal a montré que les deux niveaux excités multiplets se croisent pour une valeur élevée de la distance internucléaire R , et, dans le cas qui nous intéresse, pour $R = R_{21} \approx 10 \text{ \AA}$ (fig. 3a).

Pour de faibles valeurs de R , on a deux croisements C_1 et C_2 analogues à celui de la fig. 2a et, en C_1 , un croisement entre deux potentiels qui restent partout assez voisins. L'onde incidente donne, après passage aller-retour en C_1 et C_2 , plusieurs ondes diffusées dans chacune des deux voies ; selon le mécanisme décrit plus haut, toutes ces ondes vont se trouver mélangées en C_1 avec des différences de phase qui dépendent des temps de propagation sur les divers chemins. La différence de ces temps pour les parcours C_1C_2 et C_2C_1 ne dépend presque pas de l'angle θ de diffusion et n'est pratiquement fonction que de l'énergie E de collision. Par suite, et contrairement au cas de la fig. 2, on n'obtient pas seulement des oscillations en mesurant la section efficace différentielle d'excitation $\sigma(\theta)$ en fonction de θ (à énergie E fixée) (fig. 3b) : des oscillations subsistent également sur la courbe qui donne la section efficace totale d'excitation $\sigma_T(4'S)$ vers le niveau 4'S par exemple en fonction de E ; il en est de même pour $\sigma_T(4'S)$ (fig. 3c).

Notons cependant que l'interférence en C_1 des deux ondes, si elle est responsable des effets d'oscillation dans $\sigma_T(4'S)$, ne peut pas changer la probabilité totale de trouver un atome excité dans l'un quelconque des deux niveaux 4'S et 4'S. Par suite, la section efficace totale d'excitation $\sigma_T(4'S) + \sigma_T(4'S)$ n'oscille pas en fonction de E ; les sections efficaces totales d'excitation $\sigma_T(4'S)$ et $\sigma_T(4'S)$ ont des oscillations en opposition de phase qui disparaissent lorsqu'on fait leur somme.



Mécanisme d'excitation de Rosenthal pour l'hélium.

Les courbes de potentiel présentent ici trois croisements écartés C_1 , C_2 et C_3 (fig. a), et les interférences entre les diverses ondes donnent des oscillations dans les variations de $\sigma(\theta)$ en fonction de θ . Contrairement au cas de la fig. 2, ces oscillations ne disparaissent pas dans la section efficace totale $\sigma_T(4'S)$; les oscillations qui subsistent dépendent de la différence de temps de parcours sur les chemins C_1C_2 et C_2C_1 .

Expériences de coïncidence ion-photon

Il est donc essentiel de savoir si l'atome après collision est dans l'état 4'S ou 4'S. On pourrait penser à mesurer la perte d'énergie cinétique de l'ion He^+ , qui fournirait l'énergie de l'atome après collision (c'est la méthode classique des pertes d'énergie). Mais, comme les niveaux 4'S et 4'S ne sont séparés que par 0,008 eV, il faudrait connaître les éner-

gies cinétiques de l'ion avant et après collision avec une précision relative beaucoup trop grande pour qu'une telle méthode soit envisageable. Par contre, on peut utiliser le fait que les niveaux 4'S et 4'S sont radiatifs, et se désexcitent en émettant des photons de fréquences ν_1 et ν_2 différentes ; en mesurant cette fréquence, ce qui est aisé grâce à la

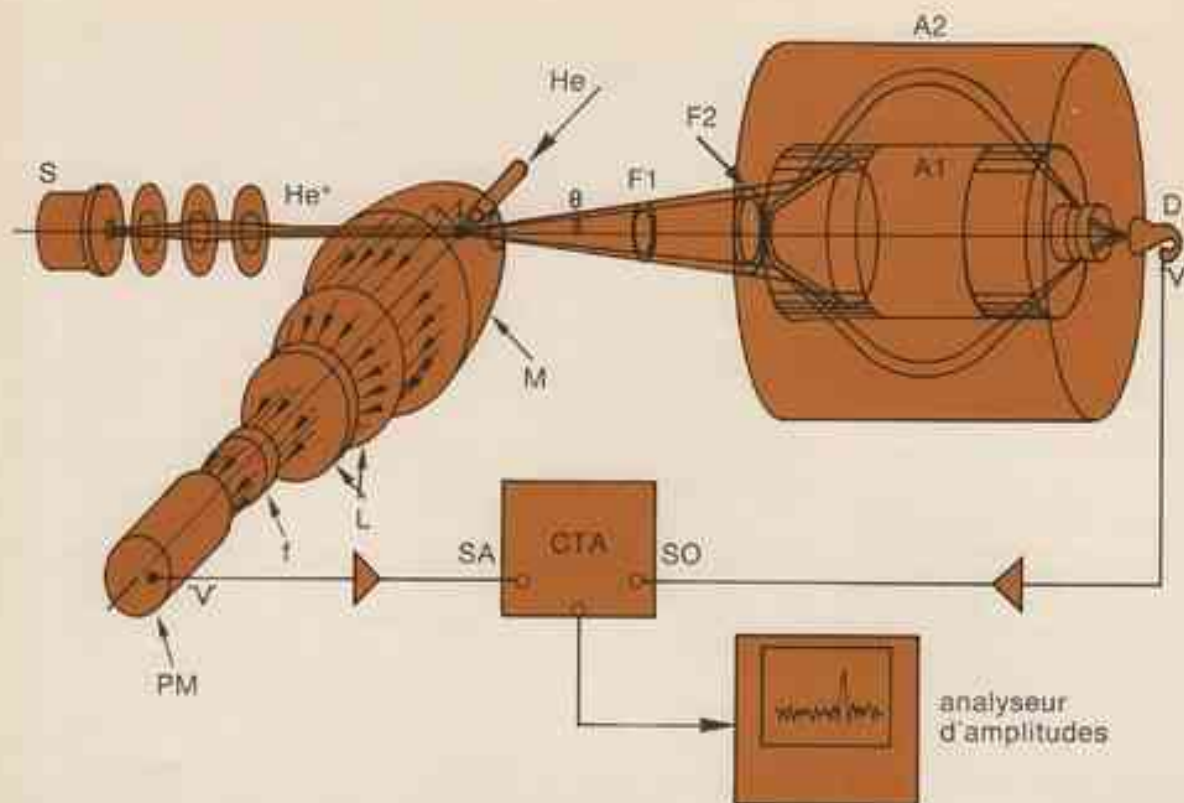
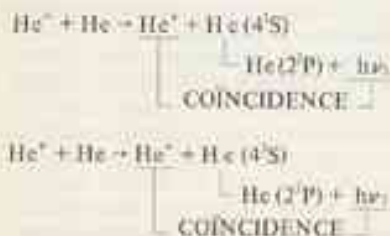


Fig. 4

Le faisceau d'ions He^+ produit par la source d'ions (S) vient croiser à angle droit un jet d'atomes d'hélium, au foyer du miroir parabolique (M). Les photons issus du volume de collision sont collectés par le miroir M et les lentilles (L), analysés par le filtre interférentiel F1 et comptés par le photomultiplicateur (PM). Les ions diffusés à l'angle θ sont prélevés par le système de fentes F1, F2, analysés en énergie par l'analyseur cylindrique (A1, A2) et enfin détectés par un multiplicateur d'électrons (D). Le convertisseur temps-amplitude (CTA) délivre des impulsions dont la hauteur est proportionnelle au temps τ qui s'écoule entre l'arrivée d'une impulsion en (SA) et l'arrivée d'une impulsion en (SO). En effectuant l'analyse en amplitude des impulsions délivrées par CTA, on obtient un "pic de coïncidences" puisque, si un ion et un photon proviennent du même événement, le délai τ est parfaitement déterminé.

grande résolution de l'analyse optique, il est possible de savoir dans quel niveau l'atome a été excité.

On emploie donc une méthode de coïncidence ion-photon : on détecte en coïncidence l'ion diffusé dans une direction donnée et le photon émis par l'atome excité, ce que l'on peut symboliser sur les schémas suivants :



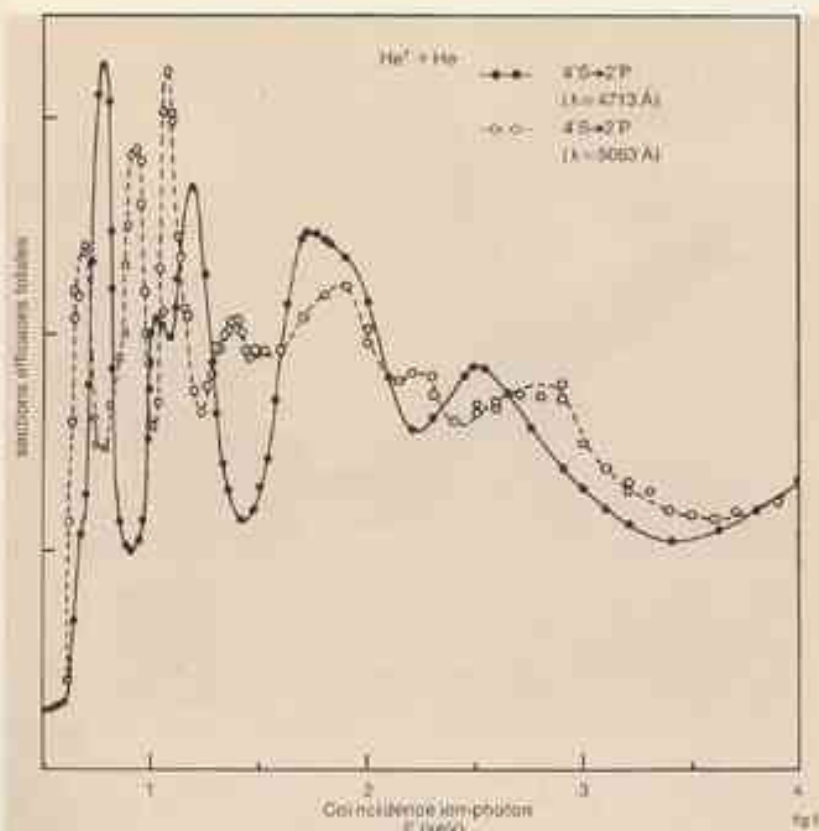
Très simple dans son principe, cette méthode est toutefois d'une mise en œuvre délicate, d'abord à cause de l'efficacité et du facteur de transmission limités des deux voies de détection (ions et photons) et surtout parce qu'il n'y a pas (ou très peu) de corrélation

entre la direction de diffusion de l'ion et la direction d'émission du photon. Parmi les photons émis dans toutes les directions, ne sont donc utiles que ceux qui correspondent à des ions diffusés dans un angle solide déterminé par l'appareillage. En outre, pour minimiser le nombre d'événements fortuits, il convient d'effectuer une analyse en énergie des ions diffusés (ne serait-ce que pour éliminer toutes les particules ayant subi un choc élastique), ce qui complique encore l'expérience.

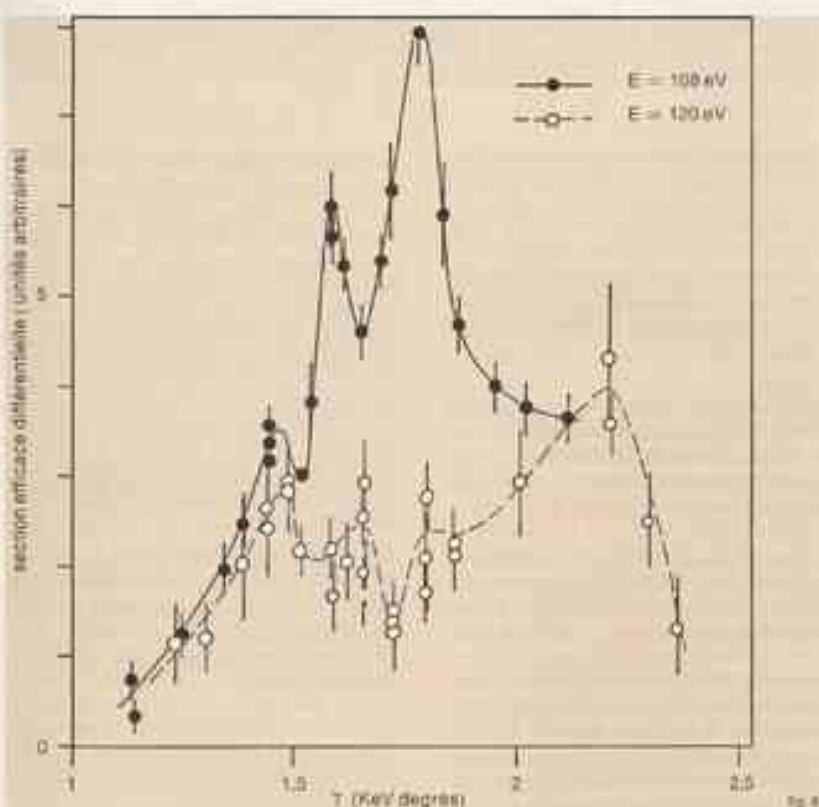
Un appareil construit sur ce principe est schématisé sur la fig. 4. Dans cet appareil, le faisceau d'ion dont on peut faire varier l'énergie de 70 à 2000 eV, croise à angle droit un jet atomique issu d'un tube multicapillaire. Les photons émis sont collectés et analysés par un filtre interférentiel suivi d'un photomultiplicateur opérant dans les conditions de détection de "photons uniques". D'autre part les ions diffusés à un angle donné sont prélevés et analysés en énergie, avec une résolution d'environ 2%.

Un appareil assez semblable, construit à la même époque aux États-Unis, fonctionne à angle fixe (11°) et énergie variable. Il semble que cette technique soit moins "directe" que la précédente car, à angle donné, le paramètre d'impact et la probabilité de transition varient avec l'énergie. L'interprétation des résultats reste donc délicate, même lorsqu'on a résolu certaines difficultés expérimentales inhérentes à la méthode, en particulier le fait que le facteur de transmission de l'appareil pour les ions dépend fortement de l'énergie. Cependant, les relations établies avec l'équipe américaine permettent une confrontation des résultats et d'entreprendre des études complémentaires ce qui est indispensable dans de telles expériences toujours très longues (4 heures par point en moyenne) et difficiles à contrôler par d'autres méthodes.

La mise au point de l'appareil a été effectuée sur l'étude de l'excitation du niveau 3^3P de l'hélium dans les collisions $\text{He}^+ - \text{He}$, de 120 à 3000 eV.



Sections efficaces totales, en valeurs relatives, pour l'excitation des niveaux $4S$ et $4S'$ de He, mesurées à partir de l'émission lumineuse qui accompagne les désexcitations $4S \rightarrow 2P$ ($\lambda = 4713 \text{ Å}$) et $4S' \rightarrow 2P$ ($\lambda = 5053 \text{ Å}$). On constate qu'à basse énergie les deux sections efficaces oscillent fortement en opposition de phase, conformément aux prévisions tirées du modèle de Rosenthal.



Sections efficaces différentielles, en valeurs relatives, pour l'excitation du niveau $4S$ de l'hélium, à des énergies de collision de 108 eV et 120 eV, auxquelles la section efficace totale est respectivement maximum et minimum. L'abscisse est l'angle réduit $T = E \cdot \theta$. À 108 eV, quatre ondes (en fait cinq dont une est négligeable) participent à la section efficace différentielle, qui présente une oscillation rapide modulée par une oscillation plus lente. Par contre, à 120 eV, il ne demeure pratiquement qu'une oscillation lente due à l'interférence de deux des ondes précédentes.

Actuellement on étudie, pour le même système, l'excitation des niveaux $4S$ et $4S'$; les résultats obtenus devraient permettre un contrôle approfondi du modèle de Rosenthal.

Dans un premier temps, on a déterminé, en détectant simplement les photons, la section efficace totale d'excitation de l'hélium sur le niveau $4S$, en fonction de l'énergie de collision. La section efficace obtenue est comparée, sur la figure 5, à la section efficace totale d'excitation de l'état $4S'$, mesurée par Dworetsky et al. On constate que les deux sections efficaces totales oscillent fortement aux faibles énergies et qu'elles sont en opposition de phase. Dès lors, les énergies les plus caractéristiques qu'il convient de choisir pour mesurer la section efficace différentielle sont celles pour lesquelles la section totale Q est extrême. Une étude plus détaillée montre que, lorsque Q est maximum, la section efficace différentielle oscille avec une amplitude modulée par un phénomène de "battements", alors que, lorsque Q est minimum les oscillations rapides disparaissent presque complètement et seule demeure une oscillation lente, en opposition avec les battements précédents. La figure 6 montre les premiers résultats obtenus sur la section efficace différentielle d'excitation du niveau $4S$ à 108 eV (Q maximum) et 120 eV (Q minimum). Ils confirment assez bien les prévisions, quoique la précision soit médiocre à 120 eV où le signal de coïncidence est très faible.

spectroscopies "beam foil" et "beam gas"

Lorsqu'on interpose, sur le trajet d'un faisceau d'ions accélérés, une cible gazeuse ou une cible solide très mince, on peut complètement changer la distribution des états de charge dans le faisceau d'ions. Par exemple, un ion Ne^+ (atome de néon auquel on a arraché un électron, encore appelé Ne II) peut donner, après traversée de la cible, des ions multichargés Ne^{2+} (Ne III), Ne^{3+} (Ne IV), etc... Si l'énergie du faisceau incident est suffisante, on peut même obtenir Ne X (ion qui n'a gardé qu'un seul électron, semblable à un atome d'hydrogène) et Ne XI (noyau sans électrons). Inversement, l'ion peut capter un électron de la cible, ce qui donne après interaction un atome neutre par exemple.

La traversée de la cible n'a pas pour seule conséquence de changer le nombre des électrons de l'ion : elle peut également provoquer une excitation des électrons qui lui restent, et l'ion émergent peut se trouver dans toute une série de niveaux excités. Lorsque ces niveaux sont radiatifs, les ions se dés-excitent spontanément dans le vide en émettant des photons. Par suite, à la sortie de la cible, le faisceau émet une lumière plus ou moins intense, qui s'affaiblit graduellement lorsqu'on s'éloigne de la cible (voir photo de tête de chapitre), dont les propriétés (longueur d'ondes, décroissance) sont caractéristiques des niveaux d'énergie excités.

L'excitation des ions par traversée de la cible (ionisation supplémentaire + passage dans un niveau excité) peut correspondre à un transfert d'énergie de plusieurs dizaines d'électron-volts, ou même plus, c'est-à-dire à des énergies importantes à l'échelle de la physique atomique ou moléculaire. De telles excitations sont difficiles, sinon impossibles, à obtenir dans toutes les sources courantes d'atomes ou d'ions excités.

Aussi n'est-il pas étonnant que l'étude de la lumière émise par le faisceau après traversée de la source ait ouvert la porte à tout un nouveau domaine de recherches en Physique atomique et moléculaire. Cette méthode, proposée en 1963 par un physicien de Manchester, L. Kay, porte le nom de "beam foil spectroscopy" quand la cible est constituée d'une feuille mince solide et de "beam gas spectroscopy" quand la cible est gazeuse. Elle a ensuite été développée à l'Université de l'Arizona par S. Bashkin, suivi ensuite par plusieurs laboratoires. Depuis 1966, elle est utilisée au Laboratoire de Spectrométrie Ionique et Moléculaire de l'Université de Lyon I.

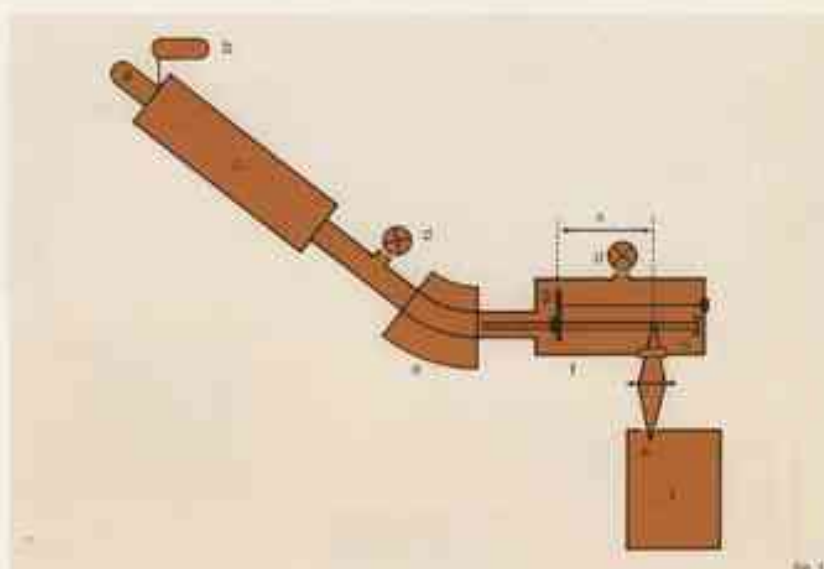


Schéma du dispositif expérimental. Dans la source haute fréquence a, le gaz qui vient du réservoir b est ionisé. Les ions extraits sont accélérés dans l'accélérateur c, puis triés par un analyseur magnétique d. Des pompes à maintien de vide maintiennent un bon vide dans l'installation et dans l'enceinte de collision f. À l'entrée de celle-ci, le faisceau d'ions traverse une très mince cible de carbone e de quelques millimètres de diamètre (la laquelle on peut imprimer un déplacement parallèle au faisceau). Les ions sont recueillis à l'aide d'un dispositif appelé cylindre de Faraday f. À la traversée de la cible, ils sont excités et la lumière émise aux différents abscisses x, recueillie par la lentille g, est étudiée au moyen du spectromètre h dans la fente d'entrée k est une focale image de la lentille.

Le dispositif expérimental

Un montage de "beam foil" (fig. 1) se compose d'une source d'ions, d'un accélérateur, d'un analyseur magnétique, d'une cible et d'un dispositif d'observation qui permet d'étudier la lumière émise par les ions.

- La source est un dispositif qui dépend essentiellement des ions que l'on veut étudier : l'ionisation des atomes est généralement obtenue par une décharge haute fréquence.
- L'accélérateur du type Van de Graaf, par exemple, est construit pour les besoins de la physique nucléaire : il peut accélérer pratiquement tous les ions.
- L'analyseur magnétique a pour but de trier les ions désirés, en fonction de leur charge Nq et de leur masse M . Il fonctionne suivant le principe de la spectrométrie de masse (un champ magnétique dévie les ions d'une façon qui dépend de N et M ; un diaphragme élimine les ions dont la déviation n'est pas celle des ions désirés).

• Laboratoire de Spectrométrie Ionique et Moléculaire, Villard-sur-Rhône (Lyon I).

• La cible, quand elle est solide, est généralement une feuille de carbone de quelques centaines d'ångströms d'épaisseur. De telles cibles sont très fragiles et le faisceau parvient fréquemment à les percer ; de ce fait, on a prévu un dispositif qui permet, de l'extérieur, de les changer. Une cible gazeuse est simplement formée par une région où, en équilibre dynamique, on entretient une certaine pression de gaz (des pompes sont bien entendu nécessaires pour éliminer le gaz qui aurait tendance à se répandre dans tout le reste de l'appareillage).

• Le dispositif d'observation est constitué d'un monochromateur à réseau et d'un détecteur lumineux très sensible (photomultiplicateur), ou encore d'un filtre interférentiel associé à un photomultiplicateur.

Un faisceau d'ions He^+ d'énergie 80 KeV pénètre dans une enceinte de cible remplie d'hélium à faible pression et traverse un film mince de carbone. La lumière bleue est émise essentiellement par les atomes d'hélium neutre de la cible gazeuse. La lumière rose est émise par les ions He^+ excités par la cible solide. On observe nettement le déclin de cette lumière à mesure que les ions s'éloignent de la cible.

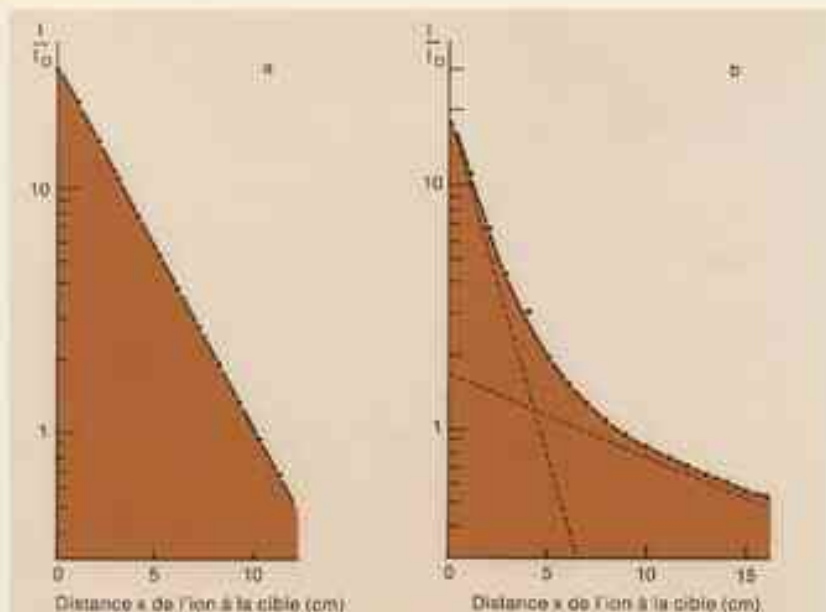


Fig. 2

Deux exemples du déclin de l'intensité des raies en fonction de la distance à la cible (en coordonnées semi-logarithmiques). En abscisses, la distance x à la cible. En ordonnées, le rapport entre I , l'intensité de la raie à cette distance, et I_0 l'intensité au niveau de la cible. L'étude porte sur deux raies émises par un faisceau d'ions carbone obtenu en alimentant la source de l'accélérateur avec de l'oxyde de carbone. En a, pour une raie de C^{2+} de 2297 Å de longueur d'onde, on peut faire passer une droite par les points expérimentaux, ce qui montre que la loi du déclin de raie correspond donc à une seule transition. La mesure de la pente de cette droite indique que la durée de vie du niveau supérieur est 7,2 nanosecondes. La courbe en trait plein de b est la courbe expérimentale qu'on obtient pour une raie de C^{2+} de 4076 Å de longueur d'onde. Cette courbe n'est pas une droite, la loi du déclin de l'intensité de cette raie n'est pas exponentielle, on observe un effet de cascade. La courbe se décompose en deux droites (tracées en pointillés) : la pente de la droite la plus inclinée correspond à la durée de vie du niveau supérieur i de la raie étudiée qui, elle-même, correspond à la transition $j \rightarrow k$. L'autre droite représente le déclin par palier d'un niveau plus élevé ($j \rightarrow i$) de durée de vie plus longue (30 nanosecondes).

Identification de niveaux excités : mesures de durées de vie

La seule identification des nombreuses raies spectrales émises par le faisceau à la sortie de la feuille est déjà très intéressante car il est possible d'en tirer des informations sur la position des niveaux d'ions dont les spectres sont mal connus. Ce travail est grandement facilité par le fait que l'analyseur magnétique utilisé dans la méthode "beam foil" permet d'obtenir une source lumineuse d'une très grande pureté chimique et même isotopique (les seules impuretés dans le faisceau sont les quelques atomes de carbone arrachés à la cible). De plus, en étudiant la dépendance de l'intensité des raies en fonction de l'énergie du faisceau d'ions incident, on peut savoir si elles appartiennent à l'atome une fois, deux fois... N fois ionisé. En tant que méthode d'identification des raies, la "beam foil" a donc des caractéristiques intéressantes qui en font un outil précieux, malgré une difficulté introduite par l'effet Doppler. En effet, les photons recueillis par le détecteur sont émis par des particules animées de vitesses élevées, et de ce fait les raies spectrales

peuvent être élargies ou déplacées. Différents dispositifs optiques permettent de réduire ou d'annuler cette perturbation.

L'intérêt le plus marquant de la méthode "beam foil", réside cependant dans l'obtention de spectres d'ions très fortement chargés, de tels spectres ne pouvant être obtenus par des sources classiques. Ainsi, avec un accélérateur de 1 MeV par nucléon, on peut observer pour l'argon les transitions de A XIII, A XIV et A XV (rappelons que l'atome d'argon neutre a 18 électrons ; il en reste donc seulement 4 à A XV). De tels états de charge ne sont observés que dans les plasmas dont la température est de plusieurs millions de degrés. Ces études sont donc essentielles en astrophysique pour l'interprétation des spectres émis par les plasmas chauds. En outre, l'étude des très courtes longueurs d'ondes émises (rayons X) donne des renseignements utiles sur les couches électroniques profondes des ions étudiés.

Une autre caractéristique intéressante de la source "beam foil" est le déclin de l'intensité émise par le jet au fur et à mesure qu'il s'éloigne de la cible. Ce déclin tient au fait que les ions excités ont une durée de vie radiative finie de

sorte que, plus l'on s'éloigne de la feuille de carbone, moins le nombre d'ions qui restent excités est grand. L'intensité lumineuse I doit donc être une fonction exponentielle décroissante de la distance x entre la région observée et la cible (fig. 2a). Comme on connaît la vitesse des ions, la mesure des variations avec x de cette intensité $I(x)$ permet d'en déduire la durée de vie du niveau supérieur de la transition. Cette méthode a été très largement utilisée et a fourni de nombreux résultats ; sa limitation principale tient aux "cascades" à partir des niveaux excités (chaque niveau est non seulement peuplé directement mais également indirectement, à partir des niveaux supérieurs d'où les ions retombent par émission spontanée) ; $I(x)$ se présente alors comme la somme de plusieurs exponentielles (fig. 2b), et les mesures de durées de vies sont beaucoup moins sûres et précises.

Étude des états multiexcités

Prenons un atome quelconque à plusieurs électrons, par exemple l'atome d'hélium qui, dans son état fondamental, a ses deux électrons dans la couche 1s. Lorsqu'on excite un des deux électrons dans les couches 2s, 2p, 3s, etc., on obtient toute une série de niveaux de l'atome dont l'énergie est limitée supérieurement par l'énergie d'ionisation E_i de l'hélium (énergie qu'il faut fournir pour arracher un électron à l'atome) ; ces niveaux excités sont ceux dont on parle habituellement. Ce ne sont toutefois pas les seuls (fig. 3) : rien n'empêche d'exciter deux électrons à la fois par exemple dans la couche 2s, ou l'un dans 2s et l'autre dans 2p, etc. ; ce sont les niveaux deux fois excités.

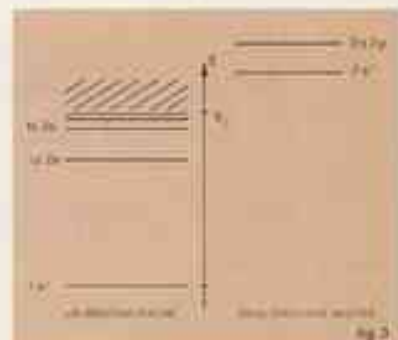


Fig. 3

Spectre des niveaux d'énergie de l'hélium. À gauche, on a représenté les niveaux où un seul électron est excité (niveaux 1s 2s, 1s 2p, etc.). L'énergie de ces niveaux est limitée supérieurement par l'énergie d'ionisation E_i (pour $E > E_i$, on a un continuum d'énergie pour l'atome ionisé, l'énergie cinétique de l'électron arraché à l'atome pourrait avoir une valeur positive quelconque). À droite figurent les niveaux multiexcités (ici deux fois excités) où les deux électrons ont quitté la couche 1s (niveaux 2s 2s, 2s 2p, etc.). L'énergie de ces niveaux étant supérieure à E_i , la plupart de ces niveaux sont autoionisants et tendent à donner un ion et un électron.

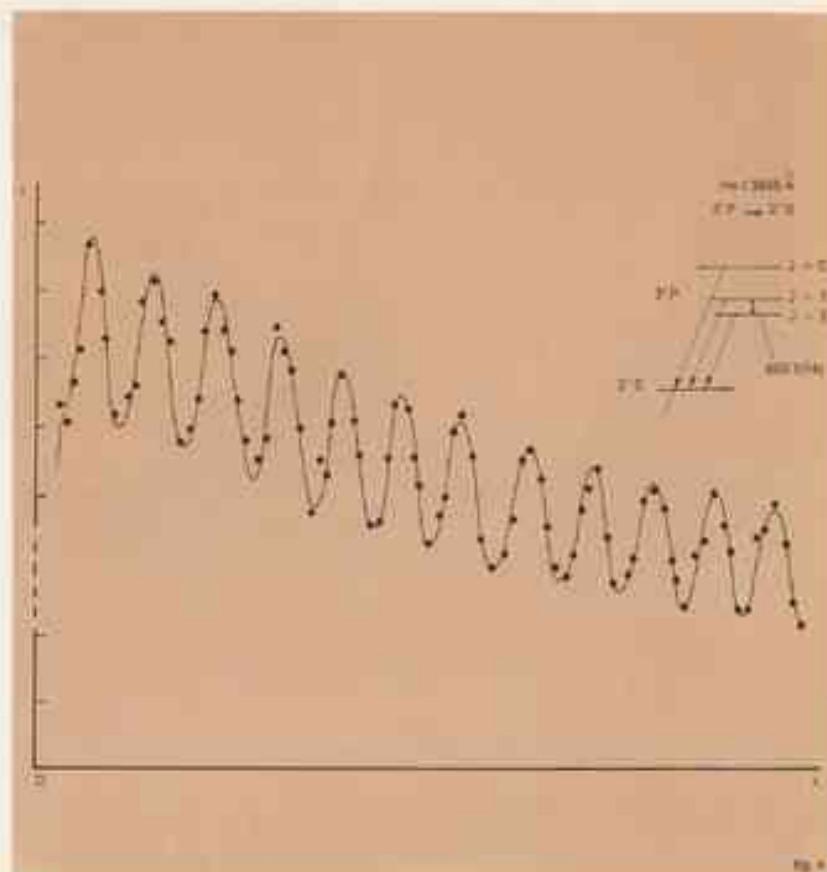
L'énergie de ces niveaux étant toujours supérieure à E_1 , il est possible énergétiquement pour l'atome de perdre un électron (il s'ionise), l'autre restant seul lié au noyau. Ce phénomène se produit pour la plupart des niveaux multi-excités, auquel cas ils sont dits "auto-ionisants". L'autoionisation est un phénomène très rapide, c'est pourquoi la durée de vie τ d'un niveau auto-ionisant est très faible (par exemple $\tau = 10^{-15}$ s, alors qu'une durée de vie radiative dans l'hélium est couramment 10.000 fois plus longue). D'après la relation d'incertitude temps-énergie, ce niveau a donc une largeur importante (inversement proportionnelle à τ), ce qui entraîne un élargissement des raies spectrales émises à partir de (ou vers) ce niveau. De tels élargissements ont été observés en beam foil spectroscopy, et plusieurs mesures de durée de vie de niveaux multi-excités ont pu être obtenues (une telle mesure n'est pas faussée, comme c'était le cas plus haut, par les phénomènes de "cascades" radiatives).

Développements récents de la méthode

Récemment, des mesures plus fines ont montré que, dans certains cas, la lumière émise par le jet peut présenter des caractéristiques plus complexes qu'un simple déclin. Ceci est dû au fait que le processus d'excitation n'est pas isotrope : le dipôle électrique associé aux ions, juste après traversée de la cible, a plus de chances d'être parallèle à la direction du faisceau que perpendiculaire (on dit que les ions excités ont un certain "alignement"). Par suite, la lumière émise par le faisceau est partiellement polarisée.

En mécanique quantique, on lie cet alignement partiel des dipôles au fait que les ions sont excités dans une "superposition cohérente" des différents sous-niveaux du niveau excité. On peut montrer que des ions excités dans ces conditions émettent une lumière dont la polarisation est modulée aux différentes fréquences de Bohr associées aux différences d'énergie entre les sous-niveaux.

Dans une expérience de "beam foil", la lumière étant émise par des particules animées de vitesses élevées, tous les processus temporels qui interviennent dans l'ion sont étalés dans l'espace (la



Déclin de l'intensité de la raie 3888 Å de HeI. Cette raie correspond à la transition $1p - 1s$ 2s. Au déclin exponentiel se superpose une modulation dont la fréquence (660 Mégahertz) correspond à la séparation des deux sous-niveaux distants de 0,0033 Å en longueur d'onde. Deux points de mesure consécutifs correspondent à un déplacement de la cible égal à 1 mm (vitesse des ions $He^+ : 4,9 \cdot 10^6$ m/s).

résolution en temps que l'on peut observer en étudiant les variations de l'intensité I en fonction de la distance x peut être meilleure que 10^{-15} s). Par suite, les courbes de déclin présenteront une ou plusieurs modulations spatiales superposées.

Effectivement, plusieurs expériences (par exemple sur les ions He^+ , Li^+) ont montré qu'il en est bien ainsi (fig. 4). En mesurant les fréquences des modulations de $I(x)$, on peut remonter aux fréquences de Bohr associées aux niveaux excités étudiés, ce qui permet d'obtenir des données sur la position de leurs sous-niveaux (structures fines et hyperfines).

D'autres expériences sont l'analogue des expériences d'"effet Hanle" décrites plus loin dans les articles relatifs au Pompage Optique (page 30) : elles consistent à appliquer le long du trajet des ions excités un champ magnétique H uniforme. Les dipôles créés lors de l'excitation sur la cible précessent alors autour de H (précession de Larmor) avec une fréquence proportionnelle à H . La mesure de cette fréquence et de l'amortissement de la modulation fournit cette

fois encore des données caractéristiques du niveau excité (facteur de Landé magnétique, durée de vie).

Cette technique de "beam foil spectroscopy" permet donc d'obtenir un grand nombre de paramètres atomiques. Dans le cas des atomes ou des ions une fois chargés, il est déjà intéressant de disposer d'une méthode assez différente des méthodes classiques. Dans le cas des ions multichargés, les techniques classiques de physique atomique utilisant des sources spectroscopiques conventionnelles sont généralement impuissantes et la méthode "beam foil spectroscopy" est la seule utilisable. Celle-ci est d'ailleurs susceptible de développements nouveaux. D'autres expériences sont possibles avec des faisceaux d'ions excités : croisements de niveaux, résonance magnétique, etc... Enfin, il n'est pas exclu que dans l'avenir on puisse substituer à la méthode, somme toute assez grossière, d'excitation par cible solide mince, une méthode plus fine d'excitation, par exemple une excitation par interaction des faisceaux d'ions avec un rayonnement laser.

les lasers à haute stabilité de fréquence

Depuis l'apparition des lasers en 1960, l'attention s'est portée sur la grande monochromaticité du rayonnement qu'ils émettent dans la gamme optique (fréquence ν de l'ordre de 10^{14} à 10^{15} Hz). Il est cependant assez vite apparu que, si la pureté spectrale "instantanée" de la fréquence émise par un laser à gaz pouvait être très élevée dans des conditions d'environnement favorables, le tableau était par contre beaucoup moins satisfaisant en ce qui concernait la reproductibilité et la stabilité de fréquence à long terme : les performances atteintes dans ces deux domaines étaient très inférieures à celles des étalons hyperfréquences tels que le jet de césium ou le maser à hydrogène.

Cette situation a été récemment modifiée par l'apparition d'une nouvelle technique, celle de l'absorption saturée, qui a déjà permis des progrès spectaculaires. Compte tenu des très nombreuses applications ainsi rendues possibles, divers laboratoires français et étrangers ont entrepris l'étude de ces nouveaux lasers stabilisés.

Stabilisation de la fréquence du laser

Les conditions d'accrochage auxquelles doit satisfaire un laser à gaz "classique" indiquent qu'un tel laser est susceptible d'osciller sur une plage de fréquences (encadrant la fréquence ν_0 de la transition utilisée) essentiellement égale à la largeur spectrale $\Delta\nu_0$, que l'effet Doppler, dû au mouvement des atomes, confère à cette transition. De plus, elles montrent que la fréquence ν de l'oscillation qui s'établit est très voisine de la fréquence d'accord ν_0 de la cavité renfermant le milieu actif. Entraînée par les décalages (mécaniques, thermiques) de celle-ci, la fréquence ν peut par suite errer sur toute la plage $\Delta\nu_0$. Il en résulte pour un tel laser une reproductibilité et une stabilité à long terme qui sont limitées à des valeurs voisines de $\Delta\nu_0/\nu_0$, soit environ 10^{-6} dans les conditions habituelles.

Pour obtenir de meilleures performances, il est nécessaire de relier plus étroitement la fréquence d'oscillation ν à la fréquence ν_0 de la transition, ce qui implique d'abord un pointé précis du centre de la raie atomique élargie par effet Doppler. Le caractère "inhomo-

gène" de cet élargissement vis-à-vis des phénomènes de saturation peut en fournir le moyen.

Considérons en effet un milieu actif soumis à deux ondes progressives opposées de même fréquence ν , l'une intense et de vecteur d'onde $\vec{k}$, l'autre d'amplitude quelconque et de vecteur d'onde $-\vec{k}$. Par suite de leur mouvement, les atomes "voient" ces deux ondes avec des fréquences différentes ν_+ et ν_- . Seuls sont en résonance avec la première les atomes pour lesquels ν_+ est égale à ν_0 , c'est-à-dire d'après l'équation classique régissant l'effet Doppler, ceux dont la vitesse v satisfait à la relation $\nu = \nu_0 + (2\pi)^{-1} \vec{k} \cdot \vec{v}$. De même, seule interagit avec la deuxième onde, la classe d'atomes différente obéissant à l'équation $\nu = \nu_0 - (2\pi)^{-1} \vec{k} \cdot \vec{v}$, c'est-à-dire de vitesse opposée à celle des atomes précédents. Il en résulte que les modifications des propriétés du milieu en cas de saturation par l'onde intense n'ont aucune influence sur le comportement de ce milieu vis-à-vis

de la deuxième onde. Le raisonnement est cependant en défaut pour la classe d'atomes correspondant à $\nu = \nu_0$ et $\vec{k} \cdot \vec{v} = 0$ (atomes voyageant perpendiculairement à la direction du rayonnement, pour lesquels l'effet Doppler est nul) qui interagit avec les deux ondes à la fois. Les propriétés du milieu vont donc présenter, en fonction de la fréquence ν , une anomalie à la fréquence ν_0 qui pourra ainsi être repérée.

Les phénomènes correspondants ont d'abord été observés dans le milieu amplificateur du laser lui-même. Il est cependant plus intéressant d'introduire dans la cavité, en série avec le milieu amplificateur, un milieu présentant une raie d'absorption de fréquence ν_0 très voisine de celle de la raie d'émission du laser (fig. 1). Le calcul de cette absorption en présence de saturation montre que la puissance de sortie d'un tel laser passe, en fonction de la fréquence d'accord ν_0 de la cavité (et donc de longueur L de celle-ci), par un maximum très aigu lorsque $\nu_0 = \nu = \nu_0$. Un dispositif d'asservissement sensible aux variations de la puissance permet alors de recentrer constamment le

« Laboratoire de l'Histoire Atomique » Orsay-Paris-Sud.

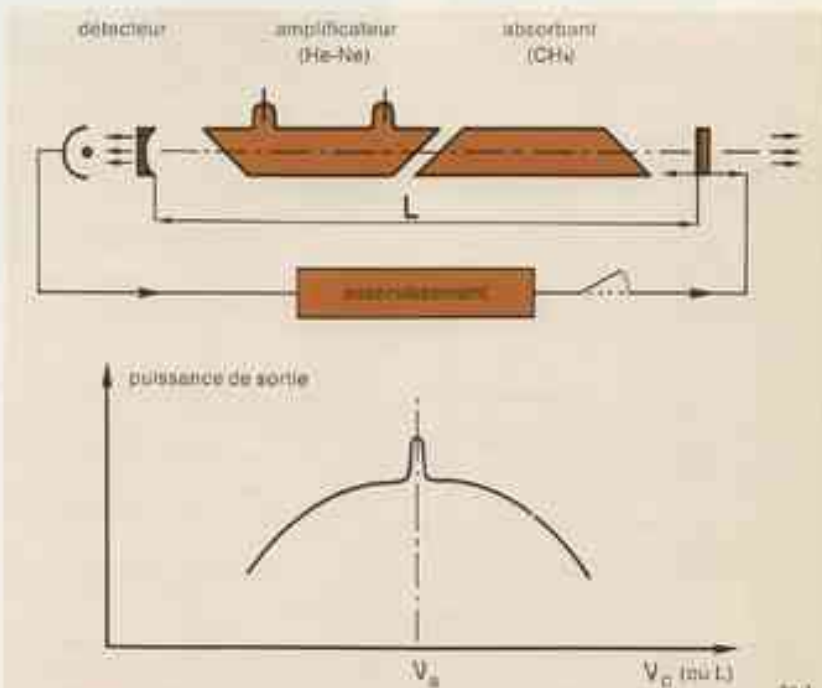


Schéma de principe d'un laser stabilisé par absorption saturée, avec la courbe de la raie d'émission du laser.

laser sur cette fréquence. Le milieu absorbant étant exempt de toutes les perturbations liées au pompage qui affectent le milieu amplificateur, il en résultera pour le laser une reproductibilité et une stabilité excellentes.

Un laser à hautes performances

Parmi les divers modèles de lasers réalisés suivant ce principe, le laser à hélium-néon stabilisé sur le pic d'absorption saturée du méthane ($\lambda = 3,39 \mu\text{m}$, $\nu = 8,84 \cdot 10^{13} \text{ Hz}$) présente actuellement les meilleures performances. La photo montre deux exemplaires d'un tel laser, installés à Orsay dans une salle climatisée sur une lourde dalle de granit les isolant des vibrations. En superposant les deux faisceaux de sortie, de fréquence ν_1 et ν_2 , sur la surface sensible d'un photodétecteur rapide, on obtient un signal de battement à la fréquence $|\nu_1 - \nu_2|$ dont les instabilités traduisent les fluctuations relatives de fréquence des deux lasers. On caractérise ces fluctuations par leur dispersion statistique σ mesurée dans des conditions bien déterminées (variance d'Allan) en fonction du temps d'observation τ (fig. 2).

La stabilité relative des deux lasers, égale à $3 \cdot 10^{-13}$ pour un temps d'observation de 100 secondes, se dégrade proportionnellement à $\tau^{-1/2}$ lorsque celui-ci diminue; pour des temps inférieurs à la constante de temps de l'asservissement, ce dernier ne corrige plus les fluctuations de fréquence et la stabilité des lasers asservis devient égale à celle des lasers libres.

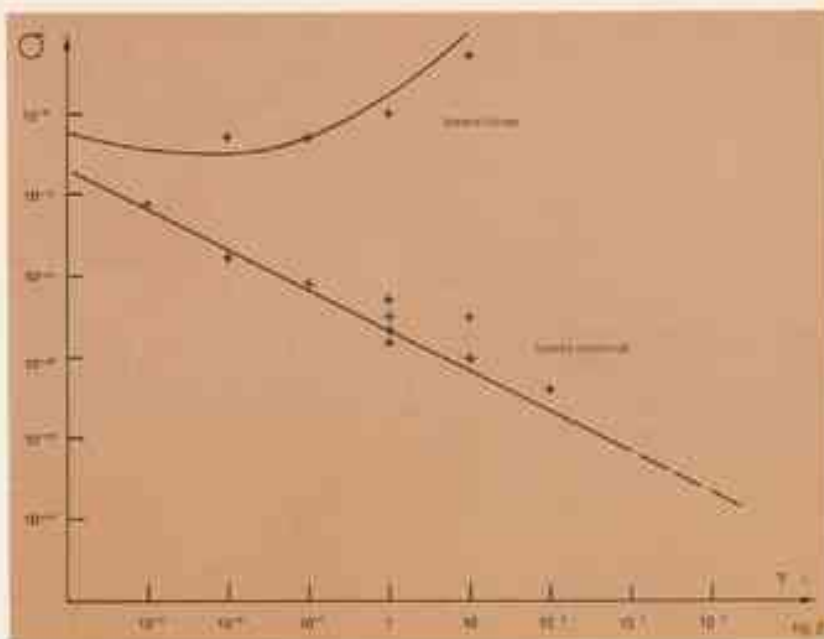
Les efforts en ce domaine portent actuellement sur l'amélioration des résultats précédents et sur la détermination de la stabilité à long terme et de la reproductibilité. Notons qu'il est du plus haut intérêt, non seulement de réaliser de tels lasers et de déterminer leur stabilité de fréquence, mais encore de mesurer d'une part la valeur absolue de cette fréquence elle-même (par comptage direct à partir de l'étalon à césium) et d'autre part la longueur d'onde correspondante (par comparaison interférométrique avec l'étalon à ^{86}Kr). Ces études difficiles, qui nécessitent la collaboration entre divers laboratoires, trouvent leur justification dans les applications qu'elles permettent d'envisager dans des domaines très variés :

- métrologie (fréquences, temps, longueurs, vitesses)
- mesure très précise de la vitesse de la lumière
- relativité expérimentale (effet de la gravitation sur la fréquence, variation éventuelle avec le temps des constantes physiques fondamentales, observation des ondes gravitationnelles)

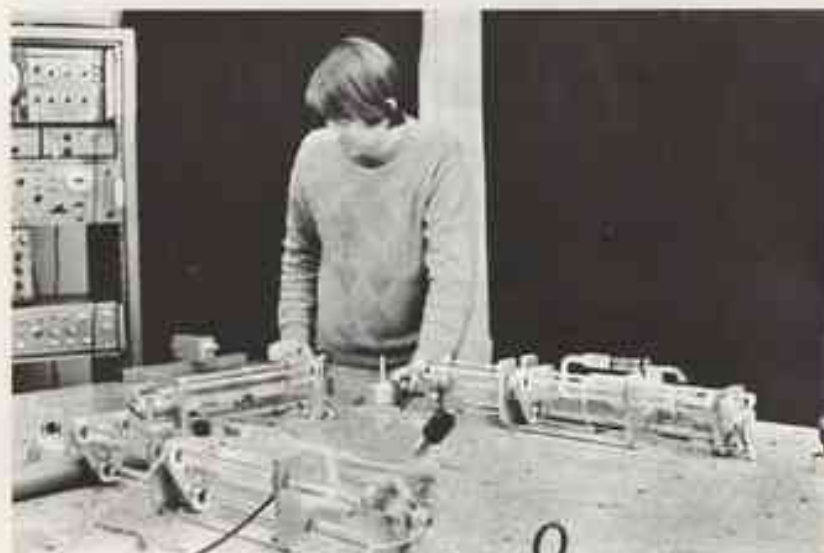
- géophysique (mesure de très petites déformations de la croûte terrestre, sismologie)

- spectroscopie (étude des très faibles déplacements de fréquence d'une raie spectrale, analyse à haute résolution). Signalons à ce propos que la mesure simultanée, au National Bureau of Standards, de la fréquence et de la longueur d'onde émises par un laser stabilisé à hélium-néon-méthane du type ci-dessus a récemment fourni pour la vitesse de la lumière dans le vide la valeur très améliorée :

$$c = 299\,792\,456,2 \pm 1,1 \text{ m/s.}$$



Stabilité relative de fréquence σ en fonction du temps d'observation τ . Les croix rendent compte des mesures effectuées au laboratoire de l'Horloge Atomique, tandis que les traits continus représentent la moyenne des résultats obtenus au National Bureau of Standards sur un dispositif similaire. On voit que ces deux déterminations indépendantes présentent un accord très satisfaisant.



Vue du montage expérimental permettant la détermination de la stabilité relative de fréquence de deux lasers stabilisés.

maser à hydrogène

Les horloges atomiques classiques (horloges à césium 133, hydrogène, rubidium 87) utilisent une transition de structure hyperfine de l'état fondamental des éléments correspondants. Cette structure hyperfine résulte du couplage entre le moment magnétique du noyau et celui de l'électron de valence. L'énergie d'interaction est faible et la fréquence de transition entre ces niveaux de structure hyperfine appartient au domaine des hyperfréquences (10^6 à 10^{10} Hertz), domaine très facilement accessible à la technique électronique. On peut démultiplier la fréquence des signaux qui permettent de déclencher la résonance atomique et obtenir une impulsion électrique, toutes les secondes par exemple, pour graduer une échelle de temps.

Les transitions du type dipolaire magnétique sont peu sensibles aux collisions qui peuvent affecter ces atomes avec ceux d'un gaz étranger ou avec la paroi (de nature convenable) d'un récipient, les interactions correspondantes étant principalement de nature électrique.

Dans le cas du césium et de l'hydrogène, il est possible d'observer la transition

dans des conditions où les déplacements de fréquence résiduels sont extrêmement faibles, et où ils peuvent être déterminés avec une grande précision : ce sont des étalons primaires de fréquence et de temps.

La transition de structure hyperfine du césium 133 sert, depuis 1965, à définir l'unité de temps, la seconde. On a par définition : $\Delta \nu_{\text{Cs}} = 9\,192\,631\,770,00 \dots$ Hz. Le maser à hydrogène possède lui aussi les qualités d'un étalon primaire de fréquence. Il présente de plus une stabilité de fréquence à court terme remarquable, ce qui lui confère un grand intérêt scientifique.

Schéma de principe

Le schéma de principe du maser à hydrogène est représenté sur la figure 1. La fréquence ν_{H} de transition entre les niveaux hyperfins E_1 et E_2 ($\nu_{\text{H}} = (E_2 - E_1)/h$) est égale à 1420 MHz environ ($\lambda = 21 \text{ cm}$). C'est cette transition qui, par ailleurs, permet aux radioastronomes de détecter l'hydrogène interstellaire.

- Laboratoire de l'Horloge Atomique - Orsay - Paris-Sud.

Un jet d'hydrogène atomique est produit à la sortie d'une source où des molécules H_2 sont dissociées dans une décharge haute fréquence. Les niveaux d'énergie E_1 et E_2 étant très voisins, ils sont sensiblement équiprobables (loi de Boltzmann). Pour obtenir un signal de résonance qui puisse être aisément détectable, il est nécessaire de produire une différence de population entre ces niveaux. Celle-ci est obtenue simplement en disposant, sur la trajectoire des atomes, un champ magnétique très inhomogène. Comme dans la célèbre expérience de Stern et Gerlach, les atomes ont des trajectoires différentes selon leur niveau d'énergie. Ici, les atomes du niveau E_1 sont collectés dans une cellule où ils peuvent séjourner pendant une seconde environ, car sa paroi interne est revêtue de téflon : la probabilité de recombinaison est particulièrement faible lors des collisions sur ce matériau et, de plus, les niveaux d'énergie sont alors peu perturbés.

Lorsque le flux d'hydrogène dans le niveau d'énergie supérieur E_2 est suffisant il se produit, par émission stimulée (effet M.A.S.E.R.), une auto-oscillation à 1420 MHz environ.

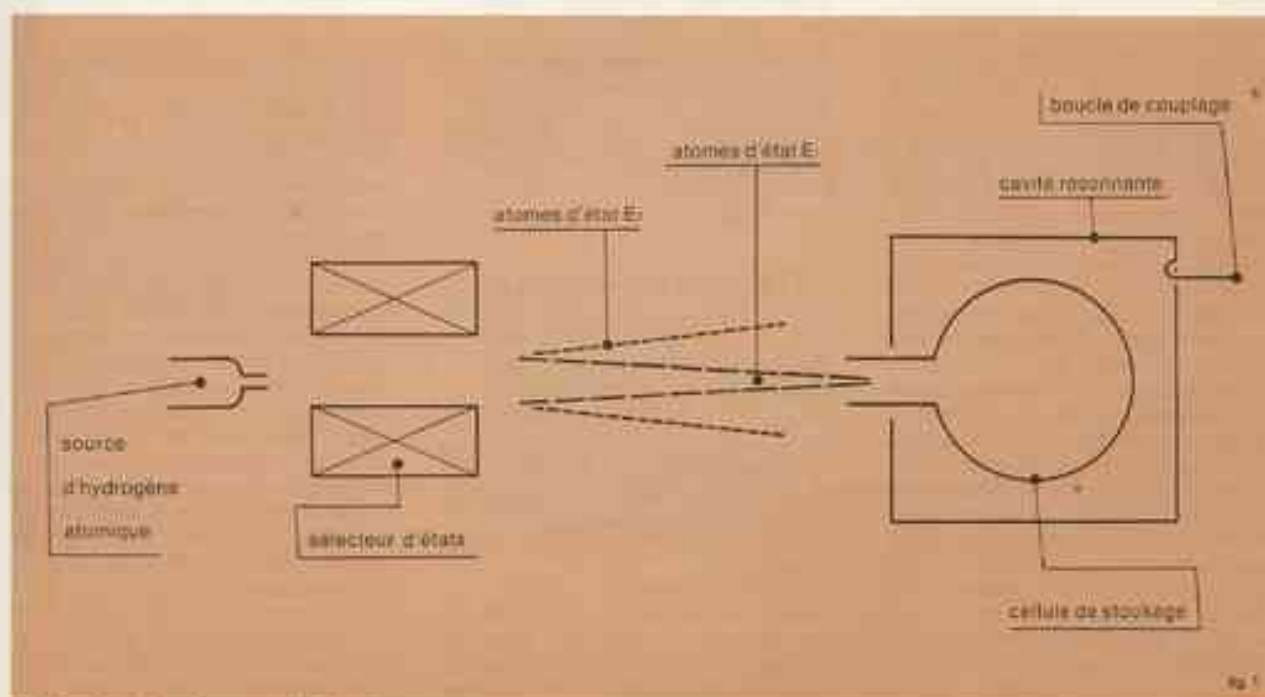


Schéma de principe du maser à hydrogène.



Photographie de deux horloges atomiques à hydrogène construites à Orsay.

Les atomes étant confinés dans un ventre du champ électromagnétique micro-onde, l'élargissement de la résonance par effet Doppler est éliminé. Les autres causes d'élargissement, par collision par exemple, sont sensiblement négligeables car la pression du gaz d'hydrogène, ainsi que celle des gaz résiduels est très faible. La largeur de la résonance atomique est alors approximativement égale à l'inverse du temps de séjour dans la cellule (relation d'incertitude), soit 1 Hz environ ou 10^{-8} en valeur relative. Le milieu atomique qui assure l'oscillation du maser étant 10^7 fois plus sélectif que la cavité résonnante qui le contient, la fréquence d'oscillation f est principalement déterminée par la fréquence de transition atomique f_0 , et il est possible d'accorder la cavité résonnante pour que l'égalité $f = f_0$ soit assurée à mieux que 10^{-11} près en valeur relative.

La photo montre deux horloges atomiques à hydrogène construites à Orsay.

La cavité résonnante est protégée des

fluctuations du champ magnétique terrestre par un blindage magnétique à six couches de mumetal. Ces horloges comprennent divers dispositifs de régulation et notamment de la température de la cavité résonnante (au millième de degré près).

Stabilité de fréquence et exactitude

Depuis 1968, la construction de deux masers à hydrogène a été entreprise au Laboratoire de l'Horloge Atomique, en vue de tester leurs performances comme horloges atomiques. Pour cela, divers systèmes de régulation ont été associés à ces masers (température de la cavité d'hydrogène, débit d'hydrogène).

Ces perfectionnements ont permis de faire fonctionner les horloges de façon continue et d'améliorer la stabilité de fréquence d'un facteur 100 environ (fig. 2).

La stabilité de fréquence à court terme est remarquable, atteignant 5×10^{-11} pour un temps d'analyse τ , compris entre 200 secondes et 1 heure. A plus long terme (sur 10 jours), la stabilité de fréquence est meilleure que 10^{-11} comme l'ont montré des comparaisons directes entre ces deux horloges. Ce chiffre a été confirmé par une comparaison à distance avec l'échelle de temps construite par le Bureau International de l'Heure, à l'Observatoire de Paris, à partir de la marche de plusieurs dizaines d'horloges à césium réparties dans le monde entier. Rappelons qu'une stabilité de fréquence de 10^{-11} correspond à un écart de marche de 1 seconde sur 3000 siècles. A titre de comparaison, la stabilité de fréquence d'une horloge à césium est de 8×10^{-12} environ pour τ compris entre 1 et 100 secondes et de 10^{-11} sur 10 jours.

De très petites corrections sont à apporter à la fréquence de transition mesurée pour passer à la fréquence de transition ν_0 dans des conditions idéalisées

(atome au repos, dans l'espace libre...). L'exactitude caractérise l'incertitude sur la fréquence ainsi extrapolée. Pour le maser à hydrogène, cette exactitude est de 2×10^{-13} (horloge à césium de laboratoire : 5×10^{-13}). Actuellement elle est limitée à cette valeur par l'imprécision sur la mesure du déplacement de fréquence résultant des collisions de l'hydrogène sur les parois de la cellule de stockage. Des progrès sont attendus dans ce domaine.

La fréquence ν_0 de transition de structure hyperfine de l'hydrogène a été mesurée en 1970, relativement à celle du ^{133}Cs , à 2×10^{-12} près. Le résultat est le suivant :

$\nu_0 = 1\,420\,405\,751,768 \pm 0,003$ Hertz. C'est la mesure la plus précise de toute la physique.

Applications spécifiques

Les applications spécifiques du maser à hydrogène résultent de sa très grande stabilité de fréquence.

L'horloge à hydrogène, plus stable que les horloges à césium, est utilisée pour contrôler la marche de ces horloges, et donc l'uniformité de l'échelle de temps qu'elles produisent. Une horloge à hydrogène, spécialisée dans cette fonction, est en cours de construction au Centre National d'Etudes des Télécommunications, avec la collaboration du Laboratoire de l'Horloge Atomique.

En radioastronomie, la résolution angulaire d'un interféromètre est d'autant plus grande que le rapport de la longueur d'onde du rayonnement à étudier à l'écartement des antennes est plus élevé. Des radio-interféromètres dont les antennes sont séparées de plusieurs milliers de kilomètres sont en service. Pour traiter les signaux reçus, il est nécessaire de les transposer en basse fréquence, sans perdre d'information sur leurs phases relatives. A cet effet, ces signaux S_1 sont mélangés avec un signal S_2 dérivé d'une horloge atomique placée auprès de chaque antenne. La cohérence de phase des signaux S_2 est assurée pendant un temps d'autant plus long que l'horloge est plus stable. C'est avec des masers à hydrogène que les meilleurs résultats sont obtenus, avec des temps de cohérence de plusieurs dizaines de minutes.

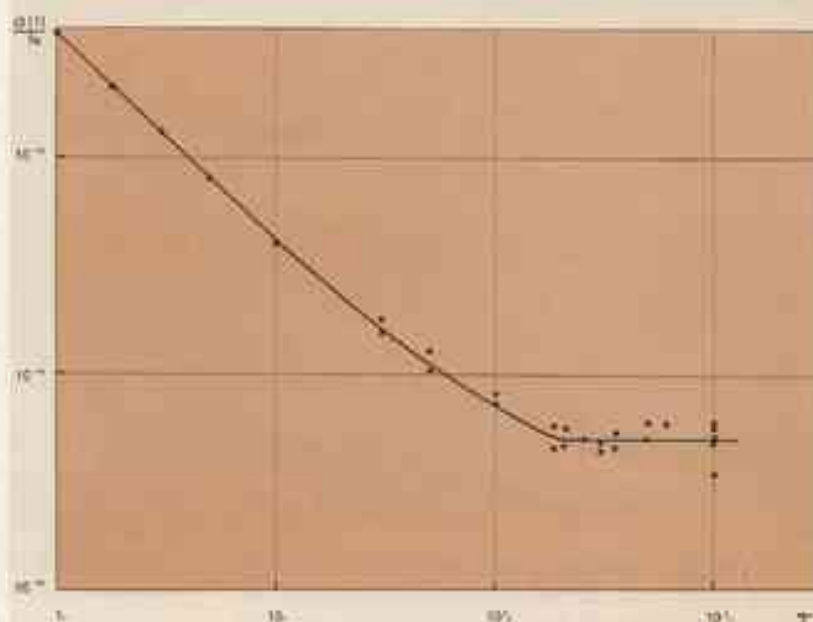


Fig. 2

Stabilité de fréquence de l'horloge atomique à hydrogène. La courbe a été établie à partir de l'analyse statistique du battiment entre les signaux produits par les 2 masers. Le diagramme de stabilité est représenté en coordonnées doublement logarithmiques. Les points correspondent aux résultats expérimentaux. Pour chaque temps de comptage τ , on détermine l'écart type $\sigma(\tau)$ des fluctuations de fréquence Δf de l'horloge, autour de sa fréquence moyenne de fonctionnement f_0 . On divise le résultat obtenu par la valeur de cette fréquence moyenne pour obtenir la stabilité en valeur relative, qui est portée en ordonnée (il serait plus correct d'appeler "instabilité", la valeur obtenue, mais l'emploi du mot "stabilité" est consacré par l'usage). Selon une méthode recommandée par le NBS, cet écart type est calculé de la façon suivante :

$$\sigma^2(f) = \frac{1}{2(m-1)} \sum_{i=1}^m (\Delta f_i - \Delta \bar{f}_{i-1})^2$$

où m est le nombre de mesures effectuées pendant l'intervalle de temps τ . Ce nombre doit être au moins égal à 100 pour que la valeur de $\sigma(f)$ soit suffisamment représentative, au sens du calcul des probabilités, du comportement de l'horloge.

L'emploi de cette méthode permet :

- de comparer les résultats obtenus dans différents laboratoires,
- d'interpréter simplement les résultats expérimentaux. C'est ainsi que la partie descendante de la courbe (pente $m = 1$) s'interprète par l'effet du bruit thermique qui s'ajoute, dans les circuits électroniques, au signal produit par le maser. Le plateau correspond à ce que l'on appelle du bruit flicque de fréquence (densité spectrale en 1/f) dont l'origine n'est pas encore très bien comprise.

Le maser à hydrogène est le générateur de fréquence le plus stable sur un intervalle de temps correspondant au temps de propagation aller et retour d'un signal radioélectrique sur la distance Terre-Lune (2,5 s), Terre-Mars (40 s) ou Terre-Jupiter (400 s). Il permet donc une mesure précise de l'éloignement d'une sonde spatiale, ainsi qu'une mesure précise par effet Doppler de sa vitesse.

Des masers à hydrogène placés aux Etats-Unis, en Australie et en Espagne ont ainsi été utilisés pendant les missions Apollo 13, 14 et Mariner.

Un test positif de la théorie de la relati-

vité générale a été réalisé en comparant le temps mis par un signal radar pour effectuer le trajet Terre-Vénus et retour, selon que la trajectoire du signal frôle ou non le Soleil. Un maser à hydrogène a été utilisé pour la mesure précise de la variation du temps de propagation.

Le maser à hydrogène est également utilisé, dans plusieurs laboratoires, pour des mesures spectroscopiques à très haute résolution (mesure de section efficace d'échange de spin, de déplacements de fréquences par échange de spin par exemple, double résonance, etc.).

le pompage optique : applications récentes

Voilà plus de vingt ans que les méthodes optiques de la Résonance Magnétique ont été mises au point. La plus connue de ces méthodes, le Pompage Optique était le fruit de la conjugaison de deux techniques antérieures : la résonance optique et la résonance magnétique. Elle a fourni de très nombreux résultats dans le domaine de la physique atomique. A l'heure actuelle, des physiciens cherchent à étendre cette méthode à de nouveaux domaines. Nous allons en décrire trois :

- les méthodes optiques de la résonance magnétique appliquées aux molécules ;
- pompage optique des électrons dans les semi-conducteurs ;
- pompage optique des centres colorés dans les cristaux.

L'idée de base commune aux expériences de pompage optique et de "Double Résonance" est la suivante : un quantum de lumière - ou photon - absorbé par la matière lui transfère non seulement son énergie mais également son moment angulaire, c'est-à-dire son état de polarisation. Inversement, lors de l'émission d'un photon par la matière, ce

photon emporte avec lui, non seulement l'énergie libérée, mais aussi une polarisation caractéristique de la transition qui lui a donné naissance. Dans un gaz, lorsqu'un photon est absorbé par un atome, celui-ci se trouve excité dans une configuration électronique correspondant au passage d'un électron d'une orbite de basse énergie, proche du noyau central, à une orbite de plus grande énergie. Si la lumière absorbée est polarisée circulairement, l'électron acquiert un moment angulaire supplémentaire dans l'état excité. Inversement, lors de la désexcitation de l'atome, l'électron cède son moment angulaire au photon émis et celui-ci est polarisé circulairement.

Que se passe-t-il maintenant si on fait agir un champ magnétique statique ou oscillant sur l'atome durant son passage dans l'état excité ? Ceci permet de changer le moment angulaire de l'atome et, par suite, la lumière qu'il réémettra changera de polarisation. On conçoit que l'on puisse de cette manière obtenir de nombreux renseignements sur les propriétés de l'atome dans l'état excité.

pompage optique des molécules

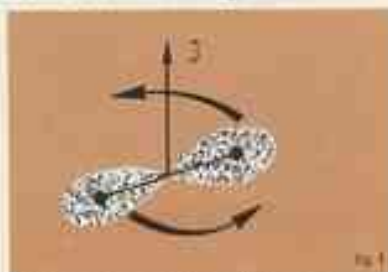
Les méthodes optiques de la résonance magnétique connaissent depuis peu un nouvel essor avec l'étude des molécules, domaine dans lequel de nombreuses interrogations restent encore sans réponses.

Jusqu'à présent, n'ont été abordées par ces méthodes que des molécules très simples constituées seulement de deux atomes liés par leur cortège électronique. Pourtant le degré de complexité d'une telle molécule, diatomique, est déjà infiniment plus élevé que celui d'un simple atome.

Les états d'énergie d'une molécule diatomique

Une molécule diatomique peut se représenter comme une petite haltère constituée de deux noyaux, entourée d'une sorte de nuage formé par l'ensemble des électrons présents. Par exemple,

une molécule d'iode I_2 est constituée de deux noyaux d'iode entourés de 106 électrons. Chaque noyau mesure environ 10^{-12} cm, la distance qui les sépare étant de l'ordre de 10^{-8} cm (1 Angström). La taille du nuage électronique est elle aussi de l'ordre de l'angström.



Rotation d'une molécule diatomique.

Laboratoire de Spectroscopie Hyperfine (Paris VI).

Dans un premier temps on peut considérer cet édifice comme rigide. Il est alors susceptible d'acquiescer une énergie de rotation E_r autour d'un axe perpendiculaire à la direction joignant les deux noyaux (fig. 1). Si J est le moment angulaire de cette molécule et I son moment d'inertie par rapport à l'axe de rotation, l'énergie ainsi acquise vaut $E_r = J^2/2I$. Elle est d'autant plus grande que le mouvement de rotation est plus rapide.

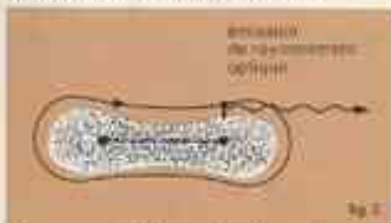
Dans un second temps, on doit tenir compte du fait que la molécule n'est pas rigide car le nuage électronique ne lie les deux noyaux que d'une façon assez élastique. En fait tout se passe comme si les noyaux étaient reliés par un petit ressort (fig. 2), l'ensemble étant susceptible d'osciller : la distance internucléaire varie alors plus ou moins sinusoidalement avec le temps. Ceci confère à la molécule une énergie de



Vibration d'une molécule diatomique.

vibration E_v , d'autant plus grande que l'amplitude de ce mouvement est grande. Le ressort que constitue le nuage électronique étant relativement raide, cette énergie de vibration E_v est environ 100 fois plus élevée que l'énergie de rotation E_r .

Enfin, comme un atome, une molécule possède des propriétés optiques dues à un ou deux de ses électrons parmi ceux qui se trouvent à la périphérie du nuage électronique : si un tel électron s'éloigne un peu du reste du nuage, il acquiert une énergie électronique E_e , qui est environ 100 fois plus élevée que l'énergie de vibration. Lorsque cet électron regagne sa place au sein du nuage, cette énergie est rayonnée sous forme de lumière visible ou ultra-violet (fig. 3).



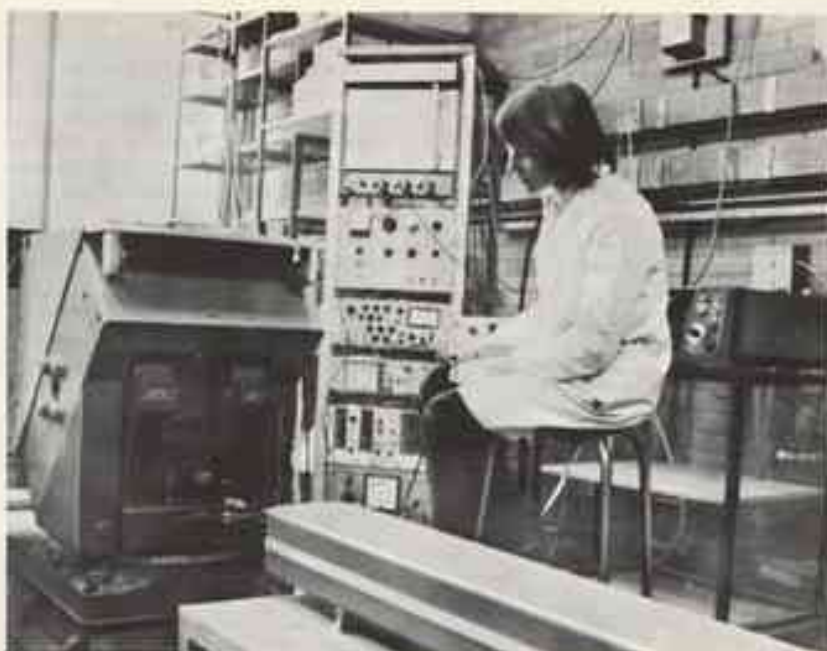
Dans une molécule excitée électroniquement, un électron s'est éloigné du reste du nuage. Cet électron peut ensuite rejoindre le nuage en émettant un photon optique.

Naturellement, ces trois types d'énergie peuvent être acquis simultanément par une même molécule qui tourne et vibre tandis qu'un de ses électrons peut s'être éloigné plus ou moins de sa trajectoire habituelle. Il convient alors de traiter globalement l'ensemble de ces problèmes à l'aide de la mécanique quantique, problème insoluble de façon rigoureuse, mais que diverses approximations permettent cependant d'aborder.

Voyons maintenant ce que sont des expériences de pompage optique appliquées aux molécules et quelles sont les informations qu'elles permettent d'obtenir.

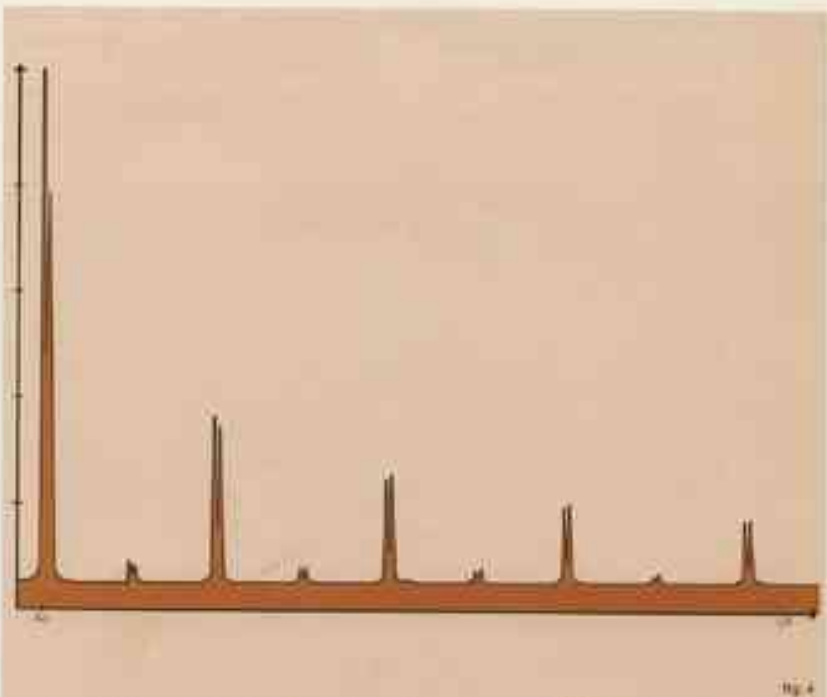
Une expérience de dépolarisation magnétique

La vapeur d'iode éclairée par de la lumière visible (le soleil par exemple) fait apparaître une intense fluorescence jaune, orange ou verte. Pour exciter cette fluorescence, on utilise par exemple une raie très monochromatique



émise par un laser à argon ou à krypton ionisé. La photo montre une vue d'ensemble du dispositif expérimental qui permet d'observer cette fluorescence sur une cellule d'iode placée entre le laser et l'électro-aimant. La fréquence ν , donc l'énergie $h\nu$ des quanta de lumière incidents est alors parfaitement

déterminée. Chaque molécule excitée à cette fréquence « possède alors des énergies de rotation, de vibration et électronique connues, ce qui se traduit par un spectre de fluorescence faisant apparaître une série de doublets que l'on peut interpréter en termes d'énergie E_r , E_v , E_e (fig. 4).



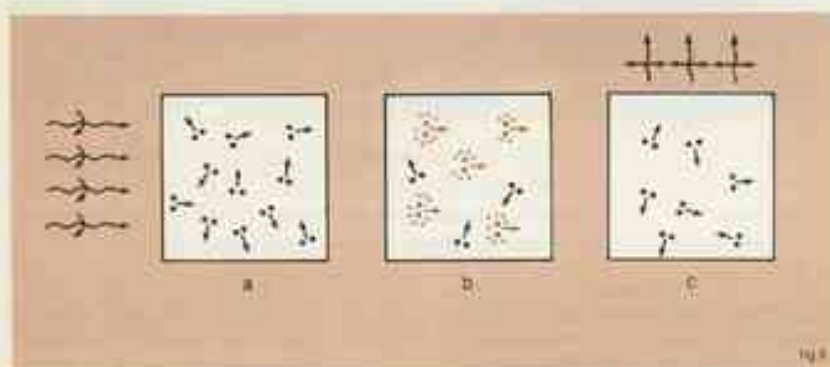
Spectre de fluorescence de l'iode excité par le rayonnement d'un laser à krypton ionisé. Les longueurs d'ondes croissent de la droite vers la gauche. λ_0 est la longueur d'onde du laser, la distance entre deux doublets successifs étant de l'ordre de 50 Å.

Il faut alors remarquer que l'on peut communiquer aux molécules, "par excitation optique", non seulement de l'énergie, mais encore du moment cinétique. Pour cela, il suffit de polariser la lumière excitatrice circulairement (lumière d'un laser après traversée d'une lame quart d'onde convenablement orientée). Le moment angulaire d'un tel faisceau lumineux, ayant une direction parfaitement définie (le long du faisceau laser), va orienter les molécules excitées de telle façon que le moment angulaire de la lumière, perdu au cours de l'absorption, soit transféré aux molécules. Le moment angulaire de rotation J des molécules, s'oriente alors de façon à satisfaire au principe de la conservation du moment angulaire. Les molécules ayant absorbé de la lumière

se trouvent donc simultanément excitées et orientées (fig. 5a et 5b). Leur orientation se traduit par une certaine polarisation de la lumière de fluorescence. Ainsi, par exemple, la lumière réémise à angle droit du faisceau d'excitation sera polarisée linéairement (fig. 5c).

L'absorption et la réémission de lumière sont séparées par un intervalle de temps τ de l'ordre de 10^{-8} sec. (durée de vie radiative du niveau excité). Pendant ce temps, il est possible d'agir de manière que les moments cinétiques des molécules perdent leur orientation ce qui se traduira par une dépolarisation de la lumière de fluorescence.

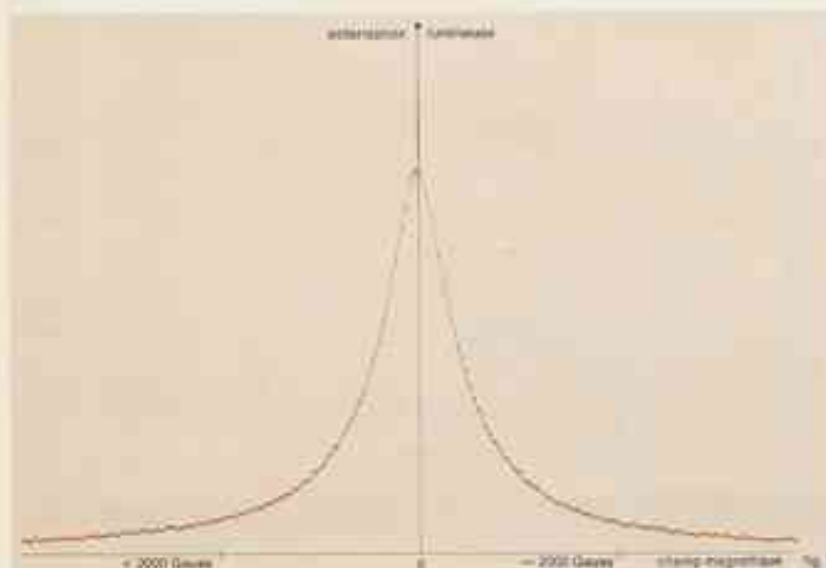
Comment produire ce phénomène ?



Les molécules, avant l'excitation optique ($\rightarrow$) ont leurs axes de rotation orientés de façon arbitraire (fig. 5a).

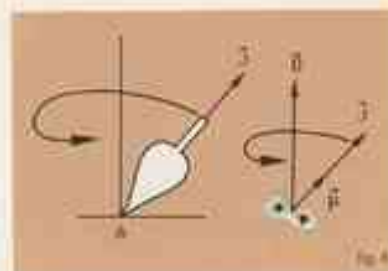
Les molécules excitées (en couleur), ayant acquis le moment angulaire du faisceau lumineux ont leurs axes de rotation orientés le long de la direction du faisceau lumineux polarisé (fig. 5b).

Lorsque ces molécules se désolent (après un temps de l'ordre de 10^{-8} secondes) en réémettant du rayonnement optique (phénomène de fluorescence), la lumière réémise à angle droit de la direction du faisceau d'excitation est polarisée linéairement (fig. 5c).



En champ magnétique nul, la lumière de fluorescence de l'ode présente un taux de polarisation élevé.

En champ magnétique élevé, le moment magnétique moléculaire ayant perdu son orientation initiale avant que la lumière de fluorescence n'ait été réémise, le taux de polarisation de la lumière de fluorescence tend vers zéro.



Sous l'action du champ de la pesanteur, une toupie dont le point de contact avec le sol A serait fixe, devrait tomber, en basculant autour de ce point. En fait, grâce à son moment angulaire J , elle ne tombe pas et précesse autour de la direction verticale du champ de pesanteur.

De la même façon, plongée dans un champ magnétique une molécule possédant un moment magnétique $\vec{\mu}$ devrait basculer afin que $\vec{\mu}$ s'aligne sur $\vec{B}$. Grâce à l'existence de son moment angulaire J , elle évite ce basculement et précesse autour du champ magnétique.

Une molécule en rotation possède, outre son moment angulaire, un moment magnétique $\vec{\mu}$ parallèle à son axe de rotation. En présence d'un champ magnétique, un moment magnétique associé comme c'est le cas ici à un moment angulaire, se met à précesser autour du champ à la manière d'une toupie précessant dans un champ de pesanteur (fig. 6). La vitesse de précession Ω est proportionnelle au champ magnétique :

$$\Omega = \gamma B$$

Si $\gamma B \ll 1$, J n'a pas le temps de précesser notablement autour de $\vec{B}$ avant que la lumière de fluorescence ne soit réémise. Celle-ci reste donc polarisée.

Au contraire si $\gamma B \gg 1$ et si $\vec{B}$ est perpendiculaire à la direction initiale de J , J précesse un grand nombre de fois autour de $\vec{B}$ avant que l'émission de lumière de fluorescence ne se produise : ce phénomène, modifiant l'orientation initiale de J , détruit également la polarisation de la lumière de fluorescence. La figure 7 présente un enregistrement typique du taux de polarisation de la lumière de fluorescence en fonction du champ magnétique. C'est une courbe dite de "dé-polarisation magnétique" de la fluorescence ou "effet Hanle".

Voyons maintenant à titre d'exemple deux types d'informations que de telles expériences peuvent nous apporter sur les molécules.

Le moment magnétique des molécules

Une molécule en rotation, nous l'avons vu, possède un moment magnétique. La figure 8 illustre trois configurations possibles mais très simplifiées du nuage électronique d'une molécule diatomique :

- celle pour laquelle il n'existe aucun moment magnétique ;

- ou un moment magnétique de même sens que $\vec{J}$;

- ou enfin un moment magnétique de sens opposé à $\vec{J}$ (fig. 8).

La réalité est en fait moins tranchée, car le nuage électronique est réparti assez largement autour des deux noyaux. En outre, la rotation des noyaux n'entraîne pas de façon rigide le nuage électronique. A la limite, on pourrait même considérer que le nuage électronique ne fait que se déformer sans réellement suivre la rotation des noyaux. Ceci montre à quel point il est difficile de prévoir théoriquement la valeur d'un moment magnétique de rotation, mais aussi à quel point ceci est important comme test de la structure électronique d'une molécule.

Les expériences précédemment décrites permettent précisément de mesurer de tels moments magnétiques. En effet, la vitesse de précession γB de l'axe de rotation est elle-même proportionnelle au moment magnétique $\vec{\mu}$. Or la courbe de la figure 7 a une largeur à mi-hauteur égale à $1/\gamma\tau$, d'où une mesure du moment magnétique.

Cette méthode ou des méthodes voisines ont d'ores et déjà permis de mesurer les moments magnétiques de nombreux états des molécules NO, I₂ et Se₂. Il semble en fait que cela permette en cause certaines structures électroniques habituellement admises.

Collisions moléculaires et dépolarisation

Si, comme on l'a déjà dit, un temps τ s'écoule entre l'excitation optique et la fluorescence, il se peut qu'en l'absence de tout champ magnétique des interactions entre les molécules elles-mêmes détruisent l'orientation créée par l'excitation optique. Qu'il s'agisse d'interactions à longue distance ou de chocs moléculaires à proprement parler, de tels phénomènes sont appelés des "collisions". Une collision entre molécules est un phénomène aux conséquences multiples. Elle peut en effet :

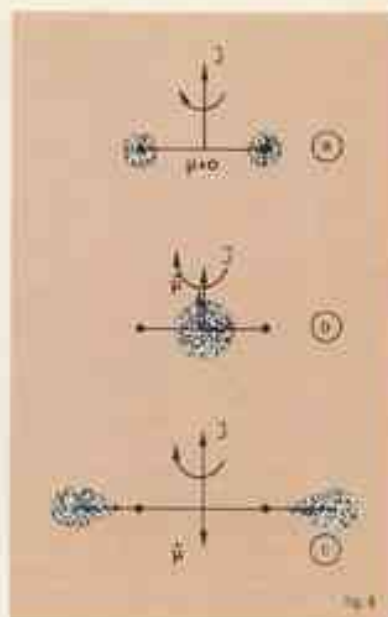
- modifier les différents états d'énergie E_0 , E_1 ou E_2 des molécules en interaction ;
- modifier l'orientation de leurs axes de rotation ;
- modifier leurs vitesses, etc...

En outre, il peut s'agir de collisions entre molécules identiques ou entre molécules différentes, ou même entre les molécules étudiées et des atomes éventuellement présents dans l'enceinte d'expérience.

S'il est facile de mesurer l'influence des collisions lorsqu'elles modifient l'énergie des molécules en présence, il est beaucoup plus difficile de déceler leur influence sur la vitesse ou sur l'orientation de ces molécules. Cependant, dans le cas considéré ici, on distingue très facilement l'influence désorientante des collisions : celle-ci se traduit par une diminution du taux de polarisation de la lumière de fluorescence. Elle se traduit également par un élargissement des courbes de dépolarisation dû au fait que la durée réelle d'interaction entre les moments magnétiques et le champ magnétique n'est plus la durée de vie τ de l'état d'énergie de la molécule, mais le temps de relaxation $T < \tau$ au bout duquel on peut considérer de toute façon la molécule comme désorientée par collision. La figure 9 montre pour l'iode la largeur des courbes de dépolarisation en fonction de la pression d'iode. L'ordonnée à l'origine est proportionnelle à $1/\tau$ puisqu'il s'agit d'une largeur extrapolée à pression nulle (en l'absence de toute collision). La pente de la droite moyenne est, elle, proportionnelle à la section efficace de dépolarisation des collisions $I_2 - I_1$.

Et l'avenir ?

Ces quelques considérations montrent que les méthodes optiques de la résonance magnétique sont applicables aux molécules et qu'elles permettent d'atteindre dans ce domaine des informations qu'aucune autre méthode ne saurait donner. Ceci ouvre la voie à de nombreuses études sur les états électroniques excités des molécules dont beaucoup de paramètres restent encore inconnus. Par ailleurs, au-delà de l'étude des molécules isolées, ces recherches conduisent à l'étude des molécules en interaction... c'est-à-dire à toute la chimie !



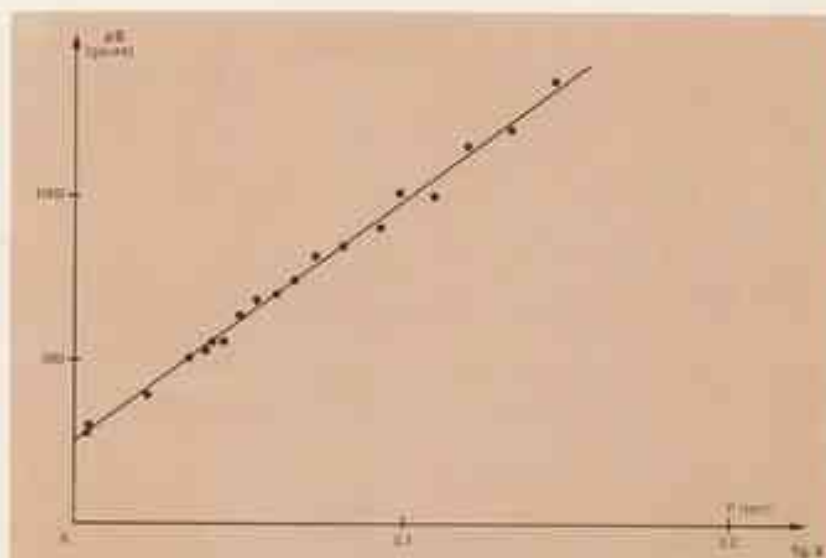
Moment magnétique d'une molécule en rotation.

Plusieurs cas sont possibles :

a) - Chaque noyau est entouré du nombre d'électrons voulus pour compenser sa charge. La rotation ne fait qu'entraîner deux édifices électriquement neutres : il n'existe alors aucun moment magnétique.

b) - Les électrons sont concentrés entre les deux noyaux sur l'axe de rotation : ils ne contribuent pas à la création d'un moment magnétique. Celui-ci est dû aux charges positives des noyaux et pointe donc dans le même sens que $\vec{J}$.

c) - Les électrons sont au-delà des noyaux. Dans le mouvement de rotation de l'ensemble, leur contribution l'emporte sur celle des noyaux : le moment magnétique résultant est alors de sens opposé à $\vec{J}$.



Élargissement des courbes de dépolarisation magnétique de I₂ avec la pression d'iode.

pompage optique des électrons de conduction dans les semi-conducteurs

Dans un semi-conducteur, les niveaux d'énergie électronique ne sont pas discrets comme dans un atome mais sont groupés en bandes d'énergie. L'état fondamental d'un semi-conducteur est celui pour lequel tous les électrons du cristal remplissent les bandes de plus basse énergie du cristal. La dernière bande remplie s'appelle la bande de valence ; les électrons qui l'occupent complètement sont assez fortement liés aux atomes qui constituent le réseau cristallin. La bande d'énergie immédiatement supérieure s'appelle la bande de conduction. Elle est séparée de la bande de valence par un intervalle d'énergie appelé bande interdite ou "gap" (fig. 1). Dans l'état fondamental, la bande de conduction est vide. L'absorption d'un photon correspond au passage d'un électron d'un état de la bande de valence à un état de la bande de conduction où les électrons, peu liés aux atomes, sont libres de se mouvoir à travers tout le cristal. L'électron manquant dans la bande de valence entraîne un déficit de charge négative et il se forme un trou porteur d'une charge positive. Électrons de la bande de conduction et trous de la bande de valence constituent les porteurs de charge responsables de la conductivité des semi-conducteurs.

Le cristal retourne à son état fondamental lorsque l'électron de la bande de

conduction retombe dans le trou de la bande de valence. Cette recombinaison électron-trou s'accompagne dans certains cas d'une émission lumineuse dont la longueur d'onde est caractéristique du semi-conducteur. C'est le phénomène de luminescence.

Pompage optique d'électrons de conduction

Lorsque le photon absorbé est polarisé circulairement, c'est-à-dire qu'il transporte un moment angulaire, ce moment angulaire est transféré aux électrons excités. Les électrons de la bande de conduction orientent alors leur moment angulaire intrinsèque, c'est-à-dire leur spin, parallèlement à la direction de propagation de la lumière. Chaque électron étant porteur d'un moment magnétique parallèle à son spin, l'orientation des électrons par absorption de photons polarisés circulairement revient à créer une aimantation globale, comme le ferait un champ magnétique. Dans des cas typiques, l'aimantation obtenue par pompage optique correspond à celle que créerait un champ magnétique de plusieurs centaines de milliers de gauss à la température de l'azote liquide.

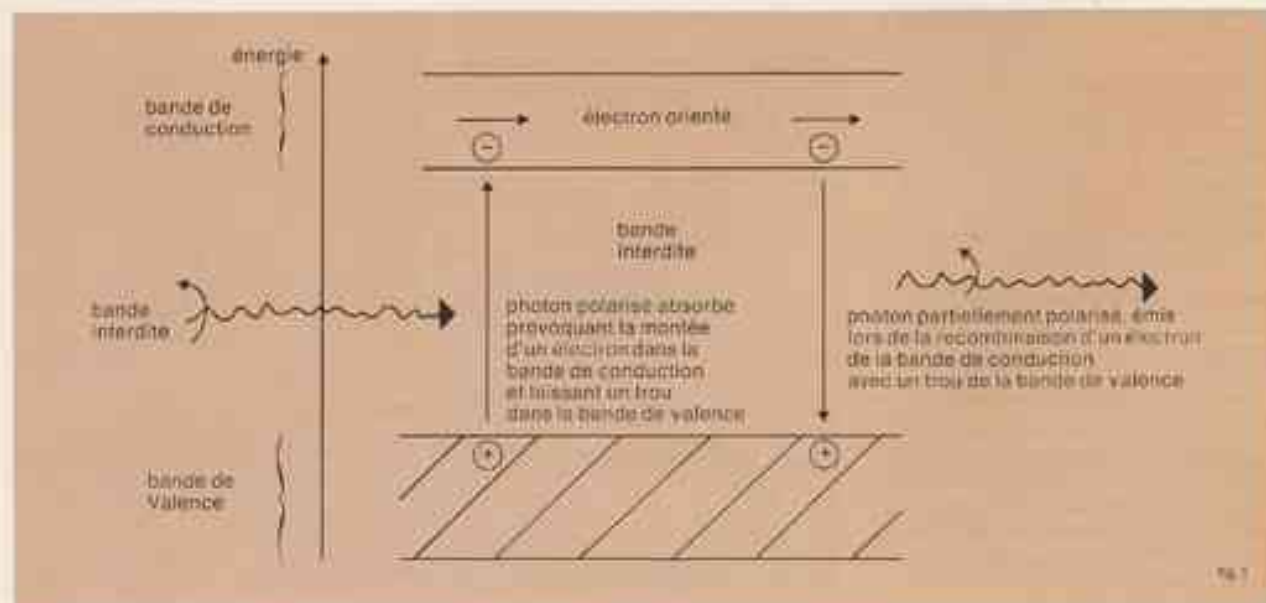
Laboratoire de Physique des Solides (Paris VII).

Inversement, lorsque des électrons orientés se recombinent, la lumière émise est polarisée circulairement, proportionnellement à leur orientation. Voyons les facteurs qui déterminent le degré d'orientation des électrons créés.

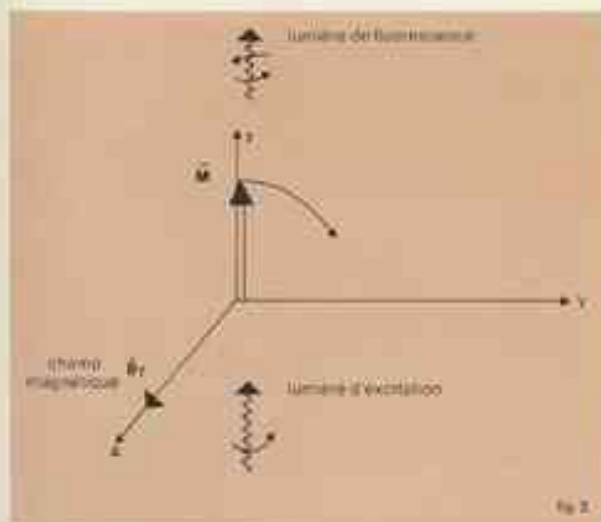
• Tout d'abord, le transfert de moment angulaire de la lumière aux électrons n'est pas total et dépend du corps considéré. Dans les semi-conducteurs courants comme l'arséniure de gallium ou l'antimoniure de gallium, l'absorption de quatre photons polarisés circulairement crée trois électrons orientés dans un sens et un électron orienté en sens contraire. L'orientation obtenue est donc :

$$\frac{3 - 1}{3 + 1} = 50 \%$$

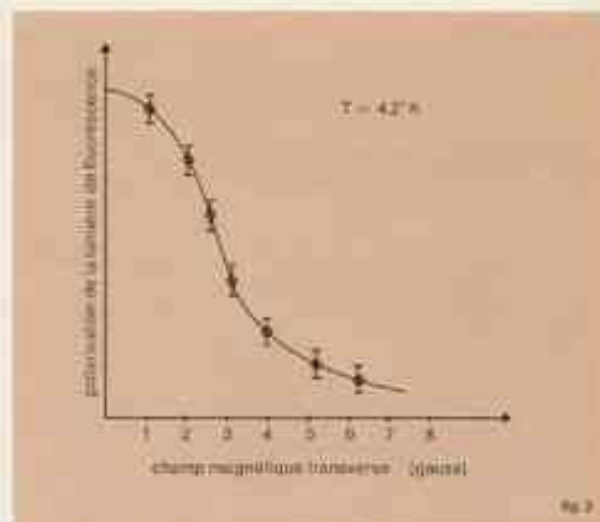
• D'autre part, les électrons mobiles soumis à des chocs sur les diverses imperfections du cristal ont tendance à se désorienter au bout d'un temps appelé temps de relaxation de spin. Ce temps est plus ou moins long suivant la pureté du cristal. Il est en général plus long à basse température qu'à haute température car l'agitation thermique des atomes constitutifs du cristal contribue à désorienter les électrons. Dans la pratique, ce temps peut varier entre une fraction de micro-seconde à la température de l'hélium liquide et une fraction de nano-seconde à plus haute température.



Structure en énergie dans un semi-conducteur.



Effet d'un champ magnétique transverse.



Dépolarisation de la lumière de fluorescence dans GaSb en présence d'un champ magnétique transverse.

• Enfin, les électrons excités par la lumière ont tendance à se recombiner avec les trous créés en même temps qu'eux au bout d'un temps moyen appelé temps de vie des électrons. Ce temps peut être très long (de l'ordre de la milliseconde dans du silicium très pur), ou au contraire très court (quelques picosecondes dans l'arséniure de gallium).

Il est clair que, si la recombinaison a lieu avant que l'électron ait eu le temps de se désorienter, les électrons qui se recombinent restent orientés et la lumière de fluorescence est partiellement polarisée circulairement. Au contraire, si la recombinaison a lieu au bout d'un temps plus long que le temps de relaxation, les électrons qui se recombinent ont perdu le souvenir de leur orientation initiale et la lumière de fluorescence n'est pas polarisée.

En résumé, la lumière de fluorescence est polarisée ou non suivant que le temps de relaxation de spin est plus long ou plus court que le temps de vie de l'électron. La mesure du taux de polarisation de la lumière de fluorescence des électrons excités par la lumière polarisée circulairement fournit donc des renseignements sur les interactions désorientantes que les électrons ont subies durant leur temps de vie.

Action d'un champ magnétique transverse

Que se passe-t-il ensuite si on applique aux électrons excités une perturbation déterminée qui a pour effet d'agir sur leur aimantation ?

Un champ magnétique $\vec{B}_1$ perpendiculaire à la direction de l'aimantation $\vec{M}$ de l'électron fait tourner cette direction

autour de celle du champ magnétique (fig. 2). La projection de l'aimantation le long de la direction de propagation de la lumière a donc diminué. La polarisation de la lumière de fluorescence, proportionnelle à la composante de l'aimantation le long de cette direction de propagation, diminue également. Dans l'antimoniure de gallium, un champ magnétique de quelques gauss suffit à dépolariser complètement la lumière de fluorescence (fig. 3). Cet effet de dépolarisation (connu depuis longtemps pour les niveaux excités des atomes d'un gaz, sous le nom d'effet Hanle) s'est montré extrêmement utile dans les expériences de pompage optique dans les semi-conducteurs. En effet, lorsqu'on mesure la polarisation de la lumière de fluorescence, il faut être certain que cette lumière est due à la recombinaison des électrons créés et non pas à un "canal" expérimental comme par exemple la lumière parasite provenant directement du dispositif d'excitation. L'application d'un champ magnétique transverse lève toute ambiguïté car ce champ serait sans effet sur la lumière parasite alors qu'il produit une dépolarisation de la lumière émise par les électrons.

Les méthodes de pompage optique des électrons permettent une application intéressante : la détection optique de la résonance paramagnétique de ces électrons. On sait que la résonance magnétique consiste à absorber de façon résonante de l'énergie en induisant des transitions entre les divers niveaux d'énergie dus à un champ magnétique (niveaux Zeeman). Ces techniques sont précieuses par le nombre de renseignements qu'elles fournissent : position et largeur des niveaux, interaction avec l'environnement. Malheureusement, les techniques habituelles d'absorption

hyperfréquence appliquées aux électrons de conduction dans les semi-conducteurs n'ont donné de résultats satisfaisants que dans un nombre limité de cas (silicium, germanium, antimoniure d'indium et quelques autres peu nombreux). La détection de la résonance, par le changement du degré de polarisation de la lumière de fluorescence au moment où l'on induit des transitions entre sous-niveaux Zeeman, a été très largement utilisée pour les niveaux atomiques des gaz. Une caractéristique essentielle de cette méthode est sa grande sensibilité due au fait que ce sont des photons optiques qui sont détectés, beaucoup plus énergétiques que les photons hyperfréquence mis en jeu dans les transitions Zeeman. Appliquée aux semi-conducteurs, cette technique a permis récemment de détecter pour la première fois la résonance paramagnétique des électrons de conduction dans l'antimoniure de gallium avec une sensibilité dix mille fois plus grande que les méthodes classiques.

Il est probablement un peu tôt pour décider de l'avenir de ces méthodes de pompage optique des électrons de conduction dans les semi-conducteurs. Cependant, on peut déjà remarquer que depuis 1968, année où elles ont été réalisées pour la première fois dans le silicium au Laboratoire de Physique de la Matière Condensée de l'École Polytechnique, elles ont connu un large développement surtout en Union Soviétique et en France. Rares sont les expériences qui n'apportent pas un aspect nouveau à ce domaine. Citons pour mémoire, outre celles que nous venons de discuter, les expériences de polarisation dynamique des noyaux, de détection optique de la résonance nucléaire, de recombinaison dépendant du spin, d'orientation des excitons...

pompage optique des centres colorés dans les cristaux

La présence d'impuretés en concentration extrêmement faible peut modifier complètement la couleur des cristaux ioniques. Cette coloration est due à des transitions optiques qui s'effectuent entre des niveaux discrets (analogues à ceux de l'atome ou de l'ion libre) se situant dans la bande interdite (fig. 1). Une coloration analogue peut être observée si on crée des défauts ponctuels dans le cristal. Un exemple classique de tels "centres colorés" est celui du centre F. Si dans un cristal de CaO (chaux vive), constitué d'ions Ca^{2+} et O^{2-} formant un réseau cubique, on crée une lacune d'ion oxygène, c'est-à-dire si on retire un de ces ions négatifs, on crée un centre positif qui peut piéger un électron (fig. 2) ; on obtient un centre appelé F^+ analogue à l'atome d'hydrogène. On peut aussi piéger deux électrons et dans ce cas on crée un centre F analogue à l'atome d'hélium. C'est le magnétisme lié à ces centres qui fait l'objet de cette étude.

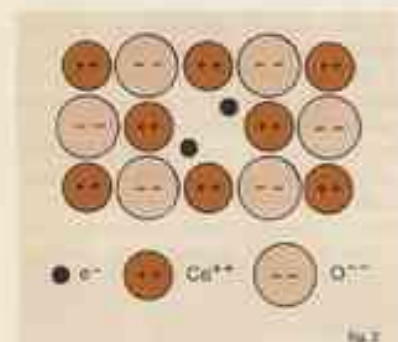


Diagramme de niveaux dans un cristal avec la bande de valence, la bande de conduction et, dans la bande interdite, les niveaux liés dus à une impureté ou à un défaut.

fig. 1

gnant d'une réorientation de l'orbite. Pour qu'elle soit possible, il faut donc qu'il existe plusieurs orbitales énergétiquement équivalentes pour l'électron. On dit qu'il y a dégénérescence orbitale. C'est le cas des fonctions p, d, f, de l'atome ou de l'ion libre. Par exemple, un électron "p" peut se trouver indifféremment dans un des trois états p_x, p_y, p_z , caractérisés par une probabilité de présence maximum sur chacun des trois axes x, y et z (fig. 3).

Ces trois états ayant la même énergie, l'électron peut sauter facilement de l'un à l'autre. C'est ce qui se produit lorsqu'on applique un champ magnétique ; s'il est dirigé selon Oz par exemple, on pourra obtenir un mouvement de l'électron dans le plan xOy par passage continu d'une orbitale x à une orbitale y et réciproquement. A ce mouvement électronique est évidemment lié un moment magnétique.

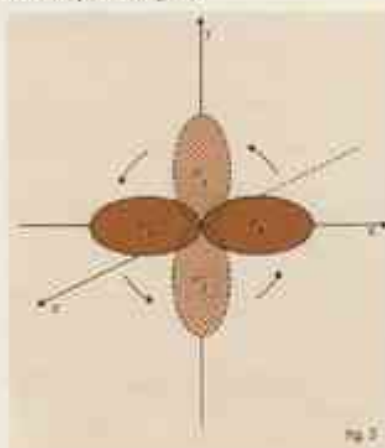
Lorsqu'on place un tel ion dans un cristal, les électrons sont soumis aux champs électriques créés par les ions voisins. Si la symétrie du site est suffisamment basse, l'équivalence des trois états orbitaux est détruite et on dit que la "dégénérescence" orbitale est "levée". Le passage d'un état orbital à un autre nécessiterait un gros apport d'énergie qui ne peut être effectué par le champ magnétique. Le mouvement orbital est alors "bloqué" et le paramagnétisme qui lui était associé disparaît. C'est le cas par exemple du niveau fondamental du chrome dans l'alumine (rubis). Lorsque la symétrie du cristal est élevée, ce blocage n'a pas forcément lieu. Par exemple, si le centre se trouve dans un cristal cubique tel que CaO, où il est entouré de six ions Ca^{2+} situés à égale distance sur des axes trirectangles, l'équivalence des trois états est conservée et on devrait s'attendre à ce que le magnétisme orbital soit peu perturbé.

Centre F dans un cristal de CaO. Deux électrons ont été piégés sur une lacune d'ion O^{2-} .

Magnétisme d'un défaut, effet Jahn Teller statique

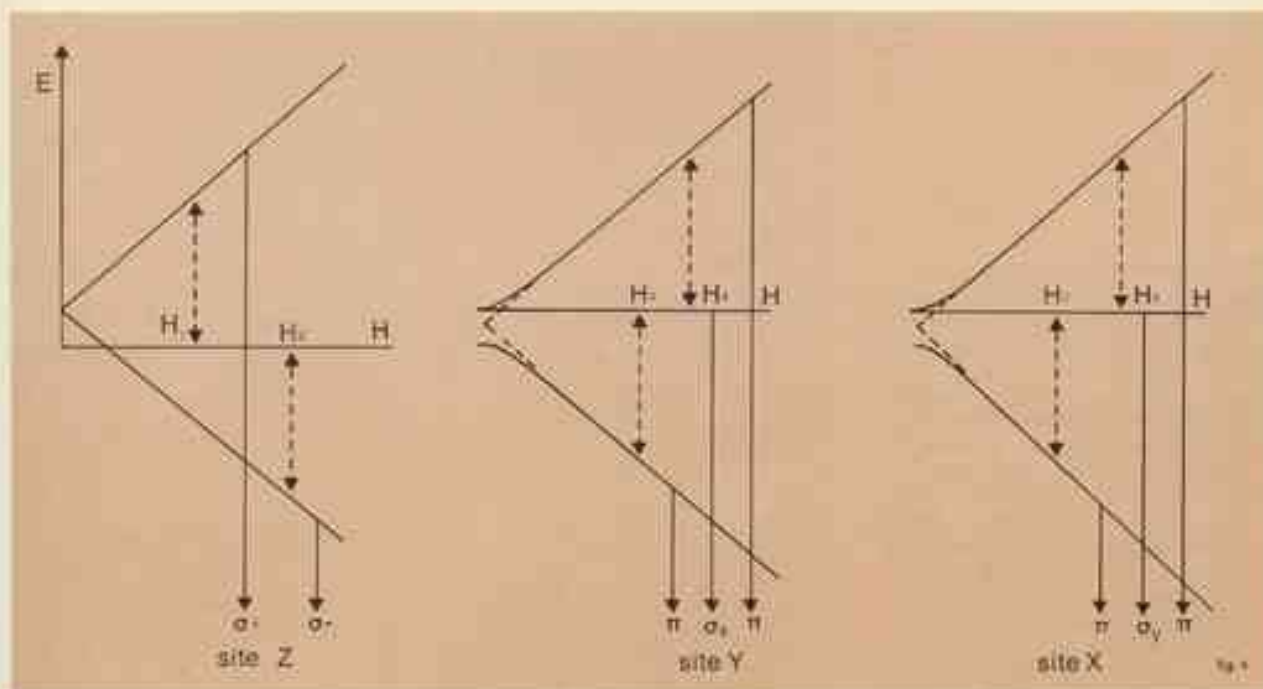
Le paramagnétisme est dû aux mouvements du ou des électrons liés aux centres F. On distingue deux types de mouvements :

- la rotation de l'électron sur lui-même, ou spin. Le fait que l'électron se trouve dans un réseau solide ne modifie pas ce mouvement puisque le réseau crée seulement des champs électriques.
- Le mouvement orbital de l'électron. Au moment angulaire est lié un moment magnétique, la réorientation de celui-ci dans un champ magnétique s'accompa-



Représentation schématisée des orbitales p_x, p_y, p_z .

Laboratoire de Spectrométrie Physique, Grenoble.



Énergie des sous-niveaux magnétiques pour chacun des trois sites X, Y, Z, en fonction du champ magnétique H appliqué suivant la direction Oz . Les flèches verticales en trait plein schématisent des transitions optiques vers le niveau fondamental 1S ; leurs polarisations lumineuses correspondantes sont indiquées en bas de la figure. Les flèches en tirets représentent les transitions induites par le champ résonant. Pour les sites Z, la résonance se produit aux valeurs H_1 et H_2 du champ, pour les sites X et Y aux valeurs H_1 et H_2 .

En fait, on constate souvent que ce magnétisme est complètement détruit. Ceci est dû à un effet prévu dès 1937 par Jahn et Teller.

Jusqu'à présent on a supposé que la position des ions voisins n'était pas modifiée par la présence de l'impureté ou du défaut. En réalité il n'en est pas ainsi car, lorsque l'électron se trouve par exemple dans l'orbitale p_z , les ions Ca^{2+} positifs situés sur l'axe Oz se rapprochent de l'origine (pour diminuer l'énergie d'interaction électrostatique). S'il se trouve dans l'orbitale p_x , ce seront au contraire les ions situés sur l'axe Ox qui auront tendance à se rapprocher. Jahn et Teller ont montré que c'est un phénomène très général et que lorsqu'un ion, dans un cristal (ou une molécule) se trouve dans un niveau d'énergie possédant une dégénérescence orbitale, le système peut toujours perdre de l'énergie en se déformant. Il en résulte que les positions d'équilibre sont obtenues pour des configurations (des noyaux) moins symétriques que celle de départ et que la dégénérescence orbitale est levée. Naturellement les trois directions x, y, z , demeurent équivalentes et dans un cristal macroscopiquement cubique, on trouvera autant de sites déformés suivant chacune de ces trois directions. Voyons maintenant un exemple d'étude d'un tel système par les méthodes de la résonance électronique.

Le centre F dans CaO

Comme on l'a déjà vu, le centre F dans le cristal de CaO est analogue à l'atome d'hélium et le diagramme de leurs niveaux est également analogue (fig. 1) : le niveau fondamental, de configuration $(1s)^2$ est 1S . La configuration excitée $(1s, 2p)$ donne lieu à deux termes : un singulet de spin 1P et un triplet 3P . La transition $^1S \rightarrow ^1P$ est autorisée et donne lieu à une bande d'absorption optique intense centrée à 4000 Å. L'énergie du niveau triplet de spin 3P est inférieure à celle du 1P . Les transitions $^3P \rightarrow ^1S$ donnent lieu à une bande d'émission centrée autour de 6000 Å. Ces transitions entre un niveau triplet de spin et un singulet sont en principe interdites, mais elles sont rendues faiblement permises par le couplage spin-orbite qui mélange les états 3P et 1P .

La durée de vie du niveau 3P est de 3 ns à 4,2 °K. C'est la résonance électronique à l'intérieur de ce niveau 3P qui a été détectée optiquement. Comme on l'a signalé précédemment, des déformations axiales de l'octaèdre de Ca^{2+} se produisent localement autour de chaque centre et tout se passe comme si les centres F étaient placés dans des sites tétraonaux. Statistiquement, les trois ensembles de sites X, Y, Z sont en nombres égaux dans un cristal cubique.

Le magnétisme orbital étant bloqué par cet effet Jahn Teller, on s'attend à n'observer que le magnétisme lié au spin $S = 1$. On sait que dans un champ magnétique H , un spin $S = 1$ peut prendre trois orientations définies par les "états propres" $|+1\rangle, |0\rangle, |-1\rangle$. Ces états sont séparés par l'énergie $g\beta H$ (β magnéton de Bohr, g facteur gyromagnétique) et des transitions peuvent être effectuées entre ces états lorsqu'on irradie le système avec une onde électromagnétique dont le quantum d'énergie $h\nu$ est égal à $g\beta H$. C'est le phénomène de résonance.

En fait, pour les centres étudiés, la dégénérescence de spin est déjà partiellement levée en champ magnétique nul (par effet du second ordre du couplage spin-orbite) et on obtient pour chacun des trois types de centres les diagrammes de niveaux indiqués sur la figure 4.

L'énergie de chacun des sous-niveaux Zeeman est portée en fonction du champ magnétique H appliqué suivant la direction Oz supposée parallèle à l'axe (001) du cristal. Le diagramme est particulièrement simple pour les sites Z dont l'axe tétraonaal est parallèle au champ H appliqué, l'énergie des niveaux $+1$ et -1 variant linéairement en fonction du champ magnétique. Pour les sites X et Y, dont les axes sont perpendiculaires à H , les diagrammes sont légèrement différents.

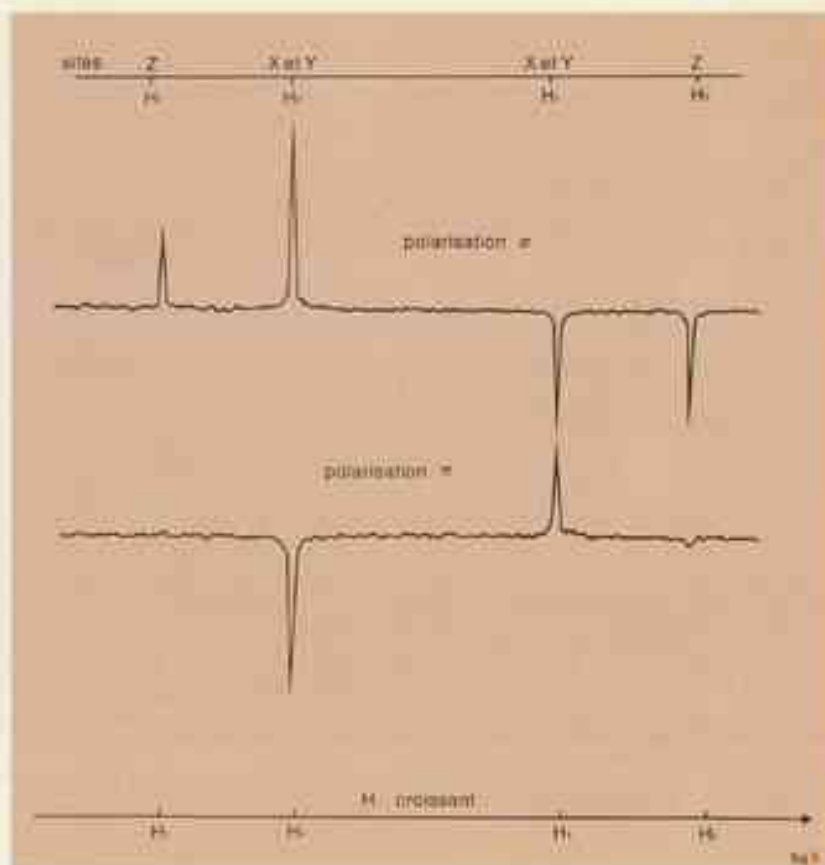
Les diagrammes d'émission à partir de chacun de ces sous-niveaux Zeeman peuvent être facilement calculés dans le cadre du modèle proposé (transitions $P \rightarrow S$ rendues permises par le couplage spin-orbite qui mélange les états P et P'). Si la direction d'observation de la lumière émise est parallèle au champ magnétique, σ_+ et σ_- sont les polarisations circulaires qui correspondent à des rotations du vecteur champ électrique respectivement dans le sens et dans le sens inverse du courant magnétisant. σ_+ , σ_- , π peuvent être définies pour des observations dans des directions perpendiculaires au champ magnétique. Elles correspondent à des polarisations linéaires dans les directions Ox , Oy , Oz . On remarque que pour les sites Z , les règles de sélection sont particulièrement simples: le niveau $|0\rangle$ n'émet pas du tout et les sous-niveaux $|+1\rangle$ et $|-1\rangle$ émettent respectivement de la lumière σ_- et de la lumière σ_+ .

Détection optique de la résonance

Comment détecter optiquement la résonance?

On place l'échantillon dans un cryostat à hélium qui permet de le maintenir à basse température (entre 4,2° K et 1,3° K). L'équilibre de Boltzmann une fois réalisé, les populations des différents sous-niveaux peuvent être notablement différentes (de l'ordre de 10 % pour des champs de 3000 Gauss). On irradie l'échantillon avec une onde hyperfréquence de fréquence fixe ν généralement choisie dans la gamme de 10.000 à 35.000 MHz. On fait alors varier le champ magnétique appliqué H . Lorsqu'on passe par la valeur H telle que la condition de résonance $h\nu = g\beta H$ est satisfaite, on induit des transitions entre les deux sous-niveaux. Si ces transitions sont induites à un rythme suffisamment rapide devant celui des processus de relaxation qui tendent à rétablir l'équilibre de Boltzmann, les populations des différents sous-niveaux Zeeman seront altérées et il en résultera des modifications de la polarisation de la lumière émise. Ce sont ces modifications que l'on détecte.

Cette méthode possède la grande sensibilité des méthodes trigger (on dé-



Variations de l'intensité lumineuse émise lors des transitions de résonance. Les sites Z résonnent pour les valeurs H_1 et H_2 du champ et X et Y pour les valeurs H_3 et H_4 . On vérifie bien les règles de sélection indiquées sur la figure 4. Lorsqu'on sature la raie $H = H_1$ des sites Z , on augmente la population du niveau $M = +1$ qui émet de la lumière σ_+ ; on attend donc une augmentation de l'intensité émise en polarisation σ_+ et pas d'effet en lumière π . C'est bien ce que l'on observe. De même, lorsqu'on sature la raie $H = H_2$ de ces mêmes sites, on diminue la population du niveau $M = -1$ et l'intensité émise en lumière σ_- diminue. Un raisonnement analogue permet d'expliquer les autres raies.

tecter un photon optique d'énergie environ 10.000 fois plus grande que le photon radiofréquence mis en jeu dans la résonance). L'observation de la résonance permet de déterminer avec une grande précision les écarts d'énergie séparant les sous-niveaux Zeeman pour des orientations variées du champ magnétique par rapport aux axes cristallins. On en déduit différents paramètres physiques tels que le facteur de Landé ou la constante de couplage spin-orbite. Par ailleurs, la nature optique de la détection de la résonance permet d'obtenir des renseignements complémentaires sur les fonctions d'ondes définissant ces sous-niveaux. La figure 5 montre les spectres de résonance obtenus en utilisant deux polarisations différentes de la lumière de fluorescence. Ces deux spectres sont nettement différents et s'expliquent très bien si on tient compte des règles de sélection indiquées sur la figure 4. Ils fournissent donc une excellente confirmation du modèle proposé.

Aspect dynamique du couplage Jahn Teller

On a supposé implicitement que les centres se trouvaient dans l'un ou l'autre des sites X , Y , Z . En fait, ces sites correspondent à des "états" à un moment donné des centres. Ces états ont approximativement la même énergie, mais ils sont séparés par des barrières de potentiel souvent élevées. Le passage de l'un à l'autre peut cependant s'effectuer par un effet tunnel analogue à celui qui donne lieu à l'inversion de la molécule d'ammoniac. Ainsi, au cours du temps un site X peut devenir Y ou Z . Les méthodes optiques, grâce à leur sensibilité et à la richesse des informations qu'elles fournissent, sont particulièrement bien adaptées à l'étude de ces phénomènes dynamiques. Différentes manifestations en ont déjà été observées (croisement de niveaux, rétrécissement des raies pour certaines orientations du champ).

condensation d'excitons en "gouttes" dans le germanium

De même qu'un électron et un proton (chargé positivement) sont liés par attraction coulombienne dans un atome d'hydrogène, de même, dans un semi-conducteur un électron et un trou libres peuvent former un exciton par un mécanisme identique.

En fait, dans les semi-conducteurs purs, l'excitation électronique d'énergie la plus basse est l'exciton et, si la température est suffisamment faible, toutes les paires électron-trou créées forment des excitons. Quand plusieurs excitons sont présents dans un cristal, leur interaction mutuelle semble être attractive, et ceci peut conduire à la formation d'un état condensé. Keldysh suggéra d'ailleurs en 1968 qu'on devait pouvoir observer dans les semi-conducteurs un phénomène de condensation d'excitons analogue à l'apparition de gouttes de liquide en présence de vapeur saturante. En effet, un bilan énergétique montre que, dans certaines conditions d'excitation et de température et dans certains semi-conducteurs, il est possible d'obtenir un système constitué de deux phases : un gaz isolant d'excitons à faible densité et des "gouttes" formées d'un plasma assez dense de paires électron-trou dont l'énergie moyenne est abaissée par rapport à celle des excitons. Il faut noter que cette phase liquide présente un caractère métallique et doit être analogue à un métal liquide tel que le sodium liquide par exemple.

Ce phénomène de condensation, qui fait actuellement l'objet de nombreuses recherches en URSS (Institut Lebedev et Ioffe) et aux Etats-Unis (Bell Laboratories, I.B.M.) est aussi étudié en France (Groupe de Physique des Solides de l'Ecole Normale Supérieure, Halle aux Vins - Paris VII).

Mise en évidence des gouttes dans le germanium

Une expérience assez simple, qui a d'ailleurs été aussi réalisée en URSS, a permis de mettre en évidence l'existence des gouttes dans le germanium. Dans cette expérience, une photodiode à jonction p-n en germanium est immergée dans un bain d'hélium liquide à 2°K et sa région n est éclairée par une

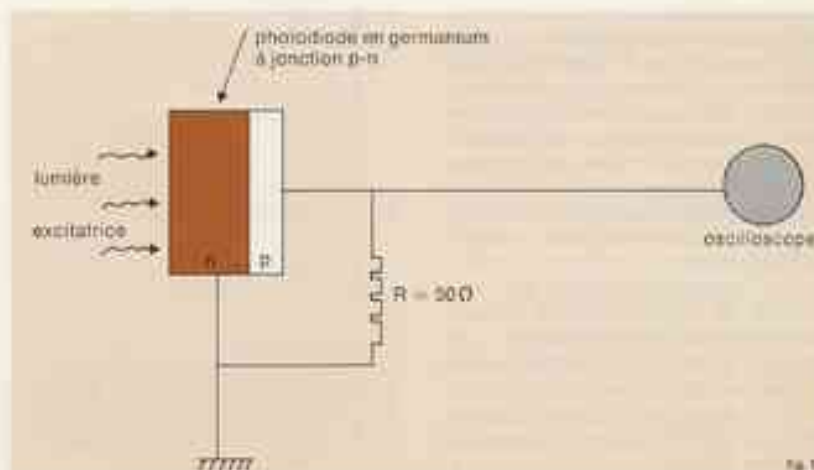
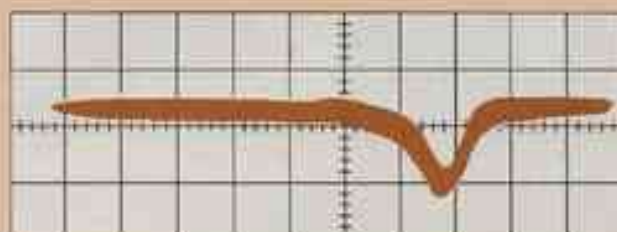


Schéma de principe du montage utilisé pour mettre en évidence l'existence des gouttes dans le germanium. Dans cette expérience, la région n d'une jonction p-n en germanium est excitée à 2° K par de la lumière continue. Pour une puissance lumineuse de l'ordre de 1 mW, on peut alors observer, sur l'écran de l'oscilloscope placé aux bornes de la résistance de 30 Ω, l'apparition de fluctuations de courant photovoltaïque dues aux gouttes de paires électron-trou.



condensation d'excitons en "gouttes" dans le germanium

Photographie de l'écran de l'oscilloscope utilisé dans l'expérience schématisée sur la Fig. 1 lorsque celui-ci est synchronisé sur les pics de courant les plus grands pouvant être détectés. Ici, le pic de courant observé correspond à une goutte comprenant $5,7 \times 10^5$ électrons. (Verticalement : 10 mV par division ; horizontalement : 10 nsec par division.)

puissance lumineuse continue de l'ordre de 1 mW (Fig. 1). Les fluctuations du courant photovoltaïque ainsi créé sont alors observées au moyen d'un oscilloscope rapide. Il apparaît dans ces conditions, sur l'écran de l'oscilloscope,

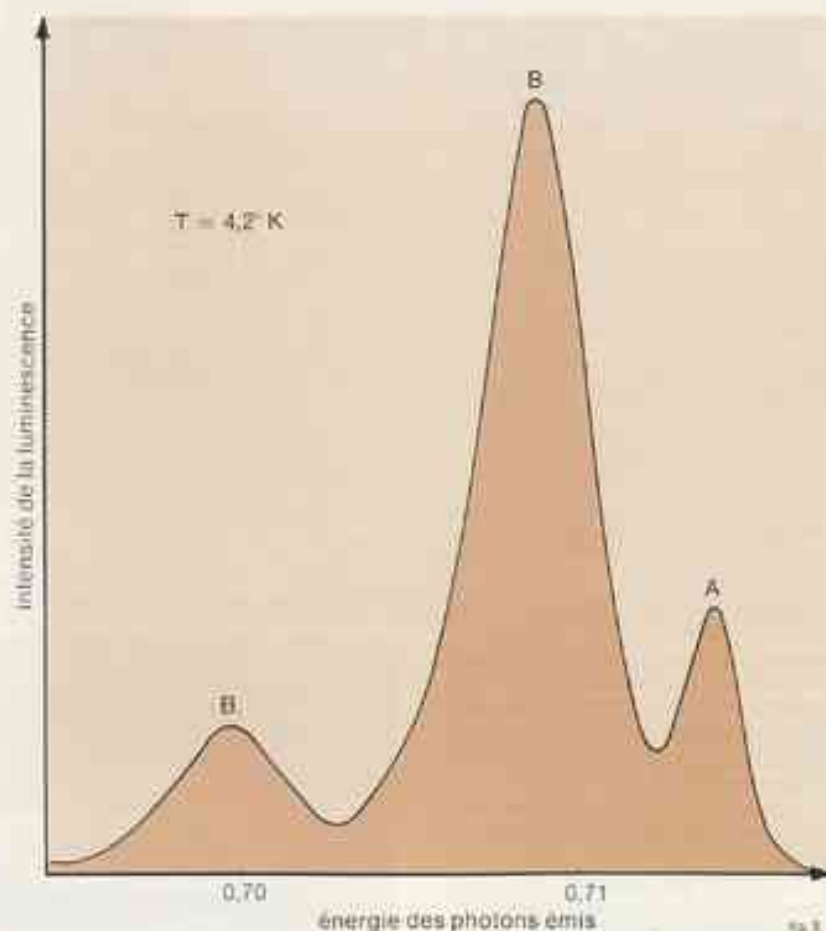
Groupe de Physique des Solides E.N.S. (Paris)

un grand nombre de pics de courant et il est même possible de synchroniser cet oscilloscope sur les plus grands d'entre eux (Fig. 2). Si on intègre ensuite le courant dans un de ces pics, on obtient directement le nombre d'électrons auquel il correspond ($3,7 \times 10^5$ dans le cas considéré ici).

On a interprété l'apparition de chacun des pics de courant observés comme étant le résultat de la destruction, à la jonction, de gouttes de paires électron-trou. En effet, si on crée de telles gouttes en éclairant la région n de la photodiode utilisée, elles vont diffuser vers la jonction proprement dite et seront dissociées par le champ électrique intense qui existe dans cette jonction. Par conséquent, si une goutte contenant N paires électron-trou est ainsi détruite, sa contribution au courant photovoltaïque doit être égale à N électrons. Si N est assez grand, un tel phénomène est certainement mesurable : c'est précisément ce qui a été observé expérimentalement puisque N peut être de l'ordre de 10^3 . Par contre, la contribution au courant photovoltaïque d'un exciton ne serait que d'un électron et cela ne pourrait pas expliquer les résultats obtenus. Cette expérience a permis également de montrer que le rayon moyen des gouttes supposées sphériques est d'environ 1μ .

D'autre part, les excitons et les gouttes correspondent à des états excités du cristal et leur désactivation peut se faire, tout au moins en partie, par émission de photons. Expérimentalement, cela se traduit par l'existence de raies de luminescence. Le spectre de la figure 3 a été obtenu à $4,2^\circ \text{K}$ en excitant un échantillon de germanium pur avec une lampe à vapeur de mercure sous pression. La raie A située à $0,7136 \text{ eV}$ est due aux excitons alors que les raies B et B₁, situées à plus faible énergie, correspondent aux gouttes. Une étude détaillée de la luminescence du germanium pur à très basse température ($T \leq 4,2^\circ \text{K}$) a permis en particulier de déterminer deux paramètres importants : la densité n de paires électron-trou dans chaque goutte et la valeur du travail de sortie correspondant à "l'évaporation" d'une paire électron-trou d'une goutte sous forme d'exciton. On a trouvé pour la densité $n = 2 \cdot 10^{17} \text{ cm}^{-3}$ et pour le travail de sortie $\Phi = 2,4 \cdot 10^{-3} \text{ eV}$. On peut remarquer que ces chiffres sont en très bon accord avec les résultats expérimentaux obtenus par d'autres groupes et aussi avec les calculs théoriques effectués en France et aux États-Unis.

Par ailleurs, des expériences de luminescence réalisées sur des échantillons



Les excitons et les gouttes peuvent disparaître en émettant des photons. En analysant l'intensité de la lumière émise en fonction de l'énergie de ces photons, on obtient un spectre de luminescence tel que celui qui est représenté sur cette figure dans le cas d'un échantillon de germanium excité à $4,2^\circ \text{K}$ par une lampe à vapeur de mercure sous pression. La raie A est due aux excitons et les raies B et B₁, qui se trouvent à plus basse énergie, correspondent aux gouttes de paires électron-trou.

soumis à une pression uniaxiale ont permis d'étudier la stabilité des gouttes et, dans ce cas aussi, les résultats sont en accord avec la théorie. D'autres expériences ont en outre fourni des informations intéressantes sur certains problèmes importants tels que le temps de formation des gouttes et leur mouvement. On a également commencé à étudier l'effet des impuretés sur les gouttes en effectuant des expériences de luminescence dans du germanium dopé à l'arsenic et à l'antimoine. Finalement, certains faits expérimentaux semblent indiquer que l'existence de

gouttes de paires électron-trou dans le silicium est vraisemblable, comme le montre d'ailleurs la théorie.

En résumé, on peut considérer que l'existence d'une transition de phase du premier ordre entre un gaz d'excitons isolant et un liquide métallique est maintenant bien établie, tout au moins dans le germanium. Il s'agit probablement là du système le plus simple présentant une transition de ce type qui, de ce fait, se prête à des calculs théoriques assez détaillés.

anomalie de Kohn et distorsion de Peierls dans des conducteurs à une dimension

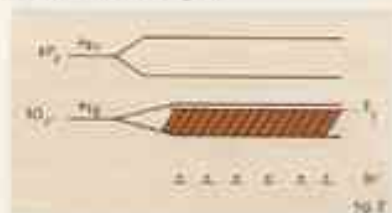
Les systèmes à une ou deux dimensions font depuis peu l'objet de l'intérêt des chercheurs qui pensent pouvoir y observer des effets caractéristiques de la dimensionnalité. Ces effets sont parfois présents dans les systèmes à trois dimensions mais sous une forme tellement atténuée que leur étude en est rendue très délicate, voire impossible.

Métal à une dimension

Les systèmes qui ont été choisis pour cette étude sont les sels de platine du type $K_2Pt(CN)_4Br_{0.36} \cdot xH_2O$, qui présentent des propriétés particulières. Ce sont des chimistes allemands qui ont attiré l'attention sur ces composés en se fondant sur leur structure cristalline. Dans la structure atomique de symétrie tétragonale apparaissent en effet des chaînes de platine parallèles à l'axe z , séparées les unes des autres par environ 10 Å et dans lesquelles les atomes de platine sont distants de 2,88 Å (à peine davantage que dans le platine métallique). Ceci a entraîné la suggestion que l'on pouvait avoir affaire à des "métaux à une dimension" provenant du recouvrement des orbitales d_z^2 des ions $Pt(CN)_4^{2-}$ (fig. 1).

À la température ambiante, la conductivité parallèlement aux chaînes de platine est environ 10^3 fois plus élevée que perpendiculairement à ces chaînes, ce qui confirme le caractère unidimensionnel des propriétés électriques. Mais

surtout, cette conductivité atteint une valeur de $10^3 (\Omega cm)^{-1}$ qui est d'un ordre de grandeur effectivement compatible avec une structure de bande électronique du type métallique. Enfin, des mesures de réflectivité optique à haute fréquence ont confirmé ce caractère métallique à la température ambiante. L'origine de ces propriétés métalliques peut se comprendre si on part de la structure de bande admise pour l'isolant $K_2Pt(CN)_6$, dans lequel la dernière bande, complètement remplie, provient précisément du recouvrement des orbitales d_z^2 du platine. L'addition d'un accepteur comme le brome, à raison d'un atome de brome pour trois atomes de platine (chaque brome prélève un électron sur les six électrons fournis par les trois platines à la bande d_z^2), peut conduire à une bande qui serait remplie au 5/6; on a alors un métal à une dimension (fig. 2).



Les ions Br^- dont les niveaux d'énergie sont plus bas que la bande d_z^2 de $K_2Pt(CN)_6$ prélèvent les électrons dans cette bande et font apparaître ainsi une bande incomplètement remplie (en couleur) dans le composé $K_2Pt(CN)_4Br_{0.36} \cdot xH_2O$.

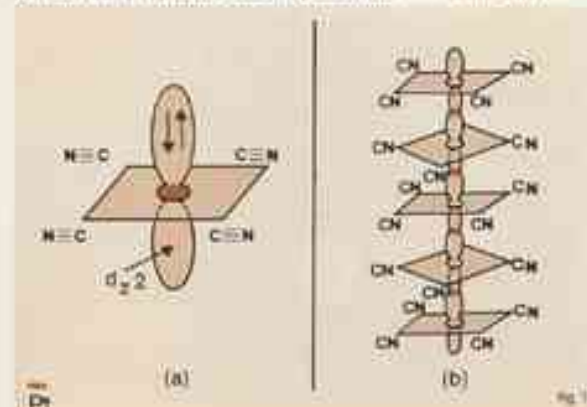
Laboratoire de Physique des Solides, Orsay (Paris Sud).

Si un modèle aussi naïf est compatible avec les résultats expérimentaux obtenus à la température ambiante, il ne peut cependant rendre compte de la variation de la conductivité électrique en fonction de la température, puisque celle-ci n'est plus que de $10^{-12} (\Omega cm)^{-1}$ à 20° K. En fait, il semble que les complications proviennent d'effets typiquement unidimensionnels.

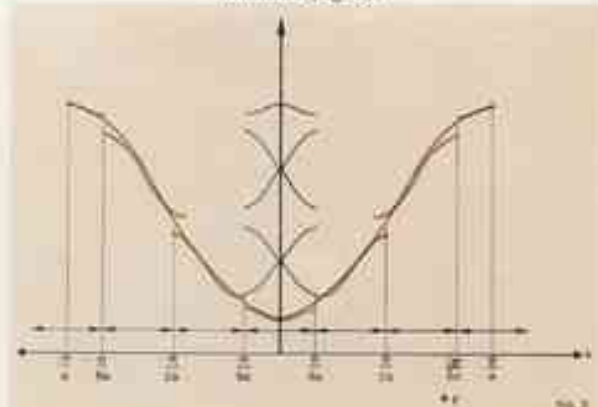
Distorsion de Peierls et anomalie de Kohn

On peut s'attendre, parmi de tels effets, à constater une "anomalie de Peierls". En effet, d'après ce dernier, à 0° K, un métal à une dimension n'est pas stable par rapport à une déformation périodique du réseau dont le vecteur d'onde q est égal à $2k_F$ (k_F , vecteur de Fermi des électrons de bande). Cette déformation périodique (qui se produit pour $q = 2 \times \frac{\pi}{a}$ dans le cas considéré),

introduit une nouvelle période (ici $\frac{\pi}{2a}$) sur le réseau, et la zone de Brillouin se trouve alors réduite de $|k| < \frac{\pi}{a}$ à $|k| < \frac{\pi}{2a}$ (fig. 3). Ainsi, pour une distorsion correspondant à $q = 2k_F$ on ouvrirait un "gap" au niveau de Fermi qui, en scindant la bande de conduction en une bande pleine et une bande vide, aboutirait à un état isolant. L'énergie de tous les états électroniques occupés et donc du système électronique entier diminue (fig. 3).



a) - Les orbitales d_z^2 et $Pt(CN)_4^{2-}$.
b) - Le recouvrement des orbitales d_z^2 parallèlement aux chaînes de platines du composé $K_2Pt(CN)_4Br_{0.36} \cdot xH_2O$.



Anomalie de Peierls pour une bande remplie au 5/6 : une déformation du réseau, de vecteur d'onde $q = 2k_F$, introduit une nouvelle période $\frac{\pi}{2a}$ sur le réseau. La zone de Brillouin se trouve réduite de $|k| < \frac{\pi}{a}$ à $|k| < \frac{\pi}{2a}$. En ouvrant un gap au niveau de Fermi on diminue l'énergie des états électroniques occupés.

Dans les systèmes à trois dimensions, le bilan global d'énergie est, en général, défavorable à une telle modification de structure (sauf peut-être pour certaines transitions métal-isolant); de ce fait, il ne subsiste qu'une faible anomalie dans le spectre de phonons : "l'anomalie de Kohn".

Dans le cas d'un système unidimensionnel, à 0°K, cette diminution d'énergie peut être plus forte que l'augmentation de toutes les autres contributions à l'énergie totale du cristal. Il s'ensuit que l'état énergétiquement le plus bas du cristal est celui avec une bande interdite au niveau de Fermi : on a alors un état isolant. Cependant, lorsque la température $k_B T$ devient de l'ordre de la bande interdite, les électrons sont excités dans la bande supérieure et le gain énergétique dû à l'ouverture de la bande interdite est perdu. Le bilan énergétique devient alors défavorable à la déformation périodique et une transition de phase cristalline se produit. Au dessus de la température critique (typiquement de l'ordre de 10 à 100°K), le souvenir de l'instabilité s'appelle anomalie de Kohn. En effet, le spectre des phonons (qui sont les déformations périodiques dynamiques), montre un minimum pour $q \approx 2k_F$ (fig. 4). A la température critique, ce minimum touche juste l'abscisse et il se perd à très haute température.

Diagramme de Rayons X

L'étude à la température ambiante, par diffusion des rayons X, du composé $K_2Pt(CN)_4Br_{0.30}H_2O$, permet d'observer un maximum de l'intensité diffusée qui, dans le modèle métallique simple décrit plus haut, est bien localisée en $\pm 2k_F$ (photo a). Le diagramme laisse voir des lignes diffuses de part et d'autre des strates qui correspondent

aux plans réciproques perpendiculaires à l'axe c^* . De telles lignes diffuses correspondent à des plans de l'espace réciproque, ce qui montre que la sur-structure est unidimensionnelle, c'est-à-dire qu'elle affecte les différentes chaînes de platine parallèles à l'axe c^* indépendamment les unes des autres. Le fait qu'on les observe de part et d'autre des strates de taches de Bragg montre (comme tous les phénomènes de taches satellites bien connus en diffraction), que la surstructure correspond à une distorsion sinusoïdale parallèle aux chaînes de platine et de période $6 \times 2,88 \text{ \AA}$. Ainsi, elle correspond à une déformation linéaire du réseau, analogue à celle que produirait un phonon longitudinal acoustique dont la longueur d'onde serait $6 \times 2,88 \text{ \AA}$.

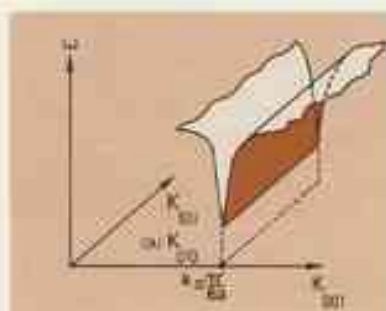


Schéma d'une anomalie dans le spectre de phonons pouvant rendre compte des diffusions en plans satellites observées aux rayons X (l'intensité diffusée par des phonons est proportionnelle à $1/q^3$).

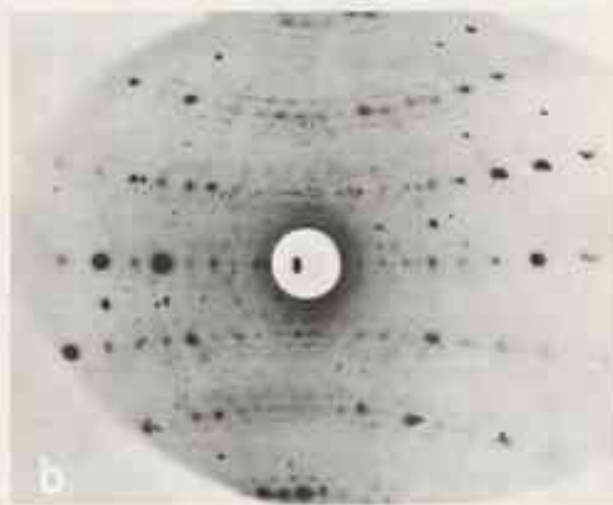
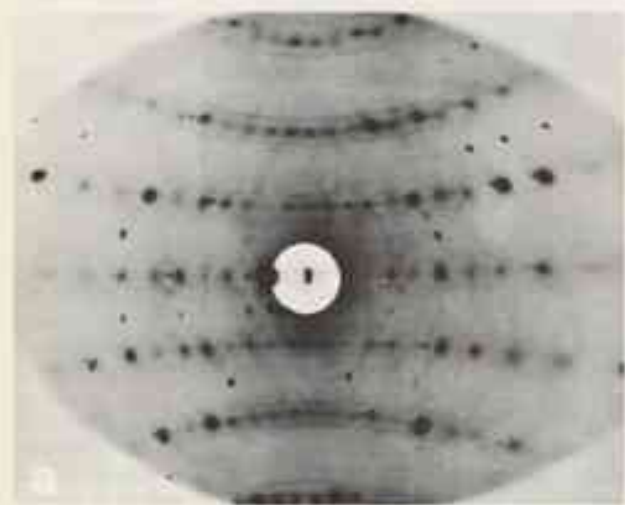
Cette diffusion très particulière subsiste sans modification sensible jusqu'à 120°K environ. Au dessous de cette température, les lignes diffuses se ponctuent (photo b) et les maxima d'intensité correspondent à des amorces de taches de Bragg qui annoncent une transition de phase vers une maille tri-di-

mensionnelle $3c \times 2a \times 2a$ si on se réfère aux paramètres de la maille stable à l'ambiante ($c = 5,76 \text{ \AA}$, $a = 9,87 \text{ \AA}$).

Les rayons X ne permettent pas en principe de distinguer une diffusion dont l'origine est une distorsion sinusoïdale statique, d'une diffusion qui provient d'un phonon longitudinal de même période correspondant à une vibration d'amplitude anormalement grande, ou ce qui revient au même dont la fréquence est anormalement faible.

Si, par contre, on tient compte des propriétés électriques, on est amené à conclure qu'à l'ambiante la diffusion observée correspond vraisemblablement à une "anomalie de Kohn géante" qui est schématisée sur la figure 4 et qui est compatible avec des propriétés métalliques. A basse température au contraire, cette anomalie s'accroît et perd son caractère uni-dimensionnel pour probablement provoquer, à une température inférieure à 77°K, la formation d'une vraie "distorsion de Peierls" statique (nouvelle structure $3c \times 2a \times 2a$) qui serait compatible avec une structure de bande de type isolant et par conséquent avec les faibles valeurs de la conductivité trouvées à 20°K.

Ceci demande bien entendu à être confirmé par des expériences effectuées à plus basse température et surtout par des expériences de diffusion inélastique des neutrons qui sont projetées à Grenoble. Mais cette conclusion est dès à présent en accord avec les calculs théoriques, et avec les premiers résultats de RMN obtenus à basse température. On observerait donc dans ces composés uni-dimensionnels une transition de phase particulière de type métal-isolant, provenant de la condensation d'un mode mou longitudinal dû à une anomalie de Kohn géante.



Les diffractions en plans satellites dans $K_2Pt(CN)_4Br_{0.30}H_2O$. a) - à l'ambiante - lignes diffuses continues. b) - à 77°K - lignes diffuses ponctuelles, montrant que la distorsion perd son caractère unidimensionnel et amorçant une transition de phase vers une maille plus grande $3c \times 2a \times 2a$.

la diffraction des rayons X par les composés magnétiques

Depuis une cinquantaine d'années, les rayons X sont utilisés pour déterminer l'arrangement des atomes dans les solides, et plus généralement dans la matière condensée. Lorsqu'un faisceau de rayons X nous tombe sur un cristal, leur longueur d'onde étant de l'ordre de grandeur des distances interatomiques ($\approx 1\text{\AA}$), il se produit un phénomène de diffraction. Les rayons X sont réfléchis dans certaines directions. La détermination de ces directions et la mesure des intensités réfléchies permettent de reconstituer la structure cristalline, c'est-à-dire l'arrangement des atomes. L'interaction entre le champ électromagnétique associé à la radiation X et les charges électriques négatives des électrons portés par ces atomes est à l'origine de cette diffusion des rayons X par les atomes.

En même temps que cette technique se développait, des physiciens étudiaient les propriétés magnétiques remarquables de certains solides et les expliquaient par un arrangement régulier de moments magnétiques portés par les atomes : ces moments peuvent être parallèles (ferromagnétisme) ou être dirigés alternativement dans deux directions opposées (antiferromagnétisme) ou encore s'arranger de façon plus compliquée, bien que périodique (ferrimagnétisme). On savait que l'on devait attribuer ces moments atomiques à un moment porté par les électrons eux-mêmes, et lié à leur spin.

Pour vérifier ces hypothèses d'arrangements périodiques des moments atomiques (ou structures magnétiques), il fallait faire diffracter sur ces substances un rayonnement de longueur d'onde encore voisine de $1\text{\AA}$ et capable d'interagir avec les spins. C'est le cas des neutrons, parce qu'ils possèdent aussi un moment magnétique. Ainsi à partir de 1949 s'est développée, parallèlement à la recherche des structures cristallines par diffraction des rayons X, une technique d'étude des structures magnétiques par diffraction de neutrons. Il y a toutefois une différence notable entre les deux techniques : alors qu'une installation de diffraction X représente un investissement à la mesure d'une équipe de recherche, un réacteur nucléaire nécessaire pour produire les neutrons est une installation lourde et coûteuse, concevable seulement dans un centre de recherche important.

Cependant, contrairement aux idées courantes exposées ci-dessus, il existe une interaction entre le spin de l'électron et le rayonnement électromagnétique, donc une possibilité de diffraction des rayons X par les structures magnétiques. On conçoit en effet assez bien que le moment magnétique de l'électron interagisse avec le champ magnétique du rayonnement. Cet effet a été calculé dès les débuts de la mécanique quantique relativiste. Sur le plan expérimental, on l'utilise pour étudier des rayons γ issus de réactions nucléai-

— Laboratoire des rayons X - Grenoble.

tes. Mais on constate qu'il a été jusqu'ici curieusement ignoré des spécialistes de la diffraction par les structures. Il y a à cela une bonne raison : lorsqu'on envoie un faisceau de rayons X sur un solide magnétique, l'intensité diffractée par les moments magnétiques de ses électrons est très faible comparée à l'intensité diffractée par leur charge électrique. Dans le cas d'un composé antiferromagnétique, pour lequel les deux types d'intensité peuvent se mesurer séparément, leur rapport est de l'ordre de 10^{-3} ; on comprend donc que ce phénomène n'ait pas été pris en considération.

Une expérience de diffraction des rayons X par un composé antiferromagnétique

Néanmoins, en 1970, deux théoriciens américains, Platzman et Tzour, en ont rappelé l'existence, évalué l'ordre de grandeur ci-dessus et estimé qu'il devrait permettre l'observation. C'est pourquoi on a tenté d'observer les intensités de rayons X diffractées par un solide antiferromagnétique.

L'observation d'intensités si faibles n'est possible qu'à deux conditions impératives :

- il faut augmenter le plus possible la puissance diffractée. On choisit pour cela un composé peu absorbant et disponible en monocristal de grande surface. On utilise l'oxyde de nickel qui présente en outre l'avantage d'être antiferromagnétique dans les conditions ambiantes. Le monochromateur que l'on doit placer dans le faisceau est en graphite orienté, matériau disponible depuis peu et qui a un très bon rendement.

- il faut réduire le plus possible les diffusions parasites qui risquent de noyer la réflexion à observer. On fait donc le vide sur le trajet du faisceau, car l'air diffuse les rayons X, et on place le compteur chargé de détecter le faisceau dans un blindage de 10 cm de plomb, pour éliminer en partie les rayons cosmiques indistinguables des rayons X.

En balayant l'angle de diffraction θ (fig. 1) autour des positions prévues

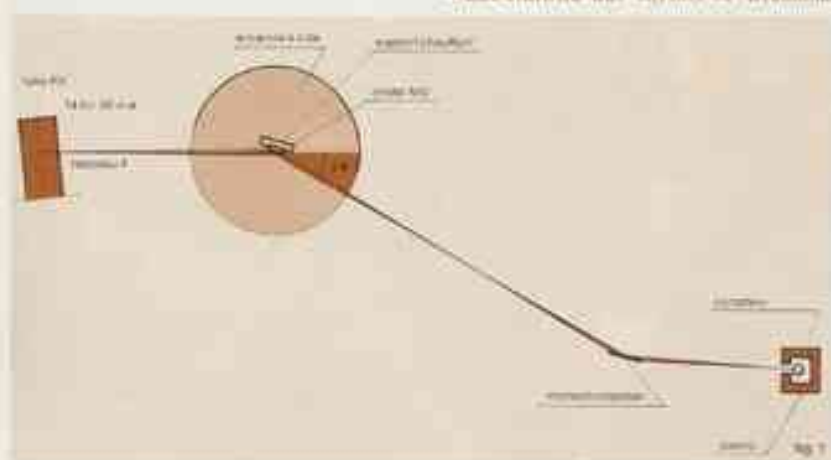
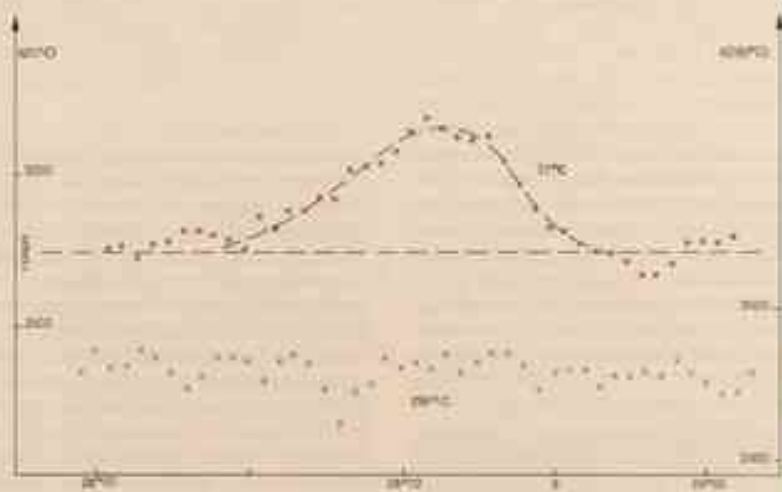


Schéma du montage utilisé. Le faisceau de Rayons X se réfléchit d'abord sur le cristal de NiO pour certaines valeurs de l'angle θ ; il subit une deuxième réflexion sur le cristal monochromateur avant de pénétrer dans le compteur. Le monochromateur élimine les fractions indésirables du rayonnement.



Intensités mesurées en fonction de l'angle de diffraction; échelle de gauche et courbe du haut, à la température ambiante; échelle de droite et courbe du bas, à 290°C. On voit apparaître à la température ambiante le pic recherché, centré sur $2\theta = 28,75^\circ$; il démontre l'existence de l'interaction des photons avec les moments magnétiques ordonnés périodiquement dans le cristal.

pour des réflexions, on a ainsi détecté des pics d'intensité ayant la hauteur attendue. Au sommet de ces pics, on compte par minute deux photons X de plus que dans le fond continu, qui présente 15 photons par minute. Trois jours de mesure sont nécessaires pour qu'un pic apparaisse clairement (fig. 2). On a vérifié que ces réflexions étaient bien dues à la structure magnétique, en chauffant le cristal à plus de 250°C : à cette température, l'ordre antiferromagnétique disparaît et on n'observe plus les pics.

On essaye actuellement d'observer le phénomène dans le cas de composés ferro ou ferrimagnétiques. Les problèmes sont ici différents, mais doivent également pouvoir être surmontés.

Deux techniques complémentaires

Les développements et les applications dépendent de l'amélioration des conditions de mesure. Il semble qu'il serait techniquement possible de gagner un facteur 10 sur les temps de mesure pour l'exemple qui a été traité. Ce ne sera pas suffisant pour que cette technique supplante la diffraction de neutrons, mais elle pourrait en devenir un complément utile, d'une part à cause de son prix, d'autre part à cause de différences entre les deux phénomènes de diffraction qui peuvent leur permettre de se compléter.

les intensités diffractées dépendent de la direction des moments magnétiques dans le solide. Cette dépendance est différente pour les deux types de diffractions; elle entraîne, dans le cas des neutrons, l'annulation de certaines réflexions visibles en principe dans une étude par rayons X. Par ailleurs, la diffraction de neutrons seule laisse parfois subsister une ambiguïté dans la détermination de la direction du moment, que les rayons X permettraient peut-être de lever. Le troisième concerne les réflexions magnétiques mesurées aux grands angles θ (fig. 1); dans le cas de la diffraction de neutrons, elles sont beaucoup plus faibles que les réflexions aux petits angles θ et donc plus difficiles à mesurer, tandis que cette décroissance est beaucoup moins importante pour les rayons X.

On peut citer trois avantages de la diffraction des rayons X. Le premier est l'étude de substances contenant un élément absorbant fortement les neutrons comme le gadolinium, qui a pourtant des propriétés magnétiques intéressantes (la diffraction de neutrons est cependant possible avec un isotope peu absorbant mais c'est une solution coûteuse). Le deuxième est lié au fait que

Indépendamment des applications possibles, ces expériences devaient être faites, car il ne faut jamais négliger la vérification expérimentale d'un phénomène prévu théoriquement, aussi tenu soit-il. De plus, la prévision théorique repose ici sur des approximations qui rendent plus nécessaire encore cette vérification.



Vue du montage utilisé : à droite le tube à Rayons X; au centre le goniomètre supportant la chambre à vide qui contient le cristal, et à gauche, abritant le compteur, le châssis de plomb à l'entrée duquel se trouve le monochromateur.

structure des alliages métalliques amorphes

Le nombre de solides que l'on peut obtenir à l'état non cristallisé a considérablement augmenté depuis que sont apparues les nouvelles techniques de trempe ultra rapide de l'état liquide, de condensation de vapeurs sur support refroidi et de dépôts par voie chimique.

Initialement les premiers solides non cristallisés connus furent les verres d'oxydes qui du point de vue structural étaient considérés comme des liquides sur-refroidis. Par contre, on a coutume de qualifier d'amorphes les éléments ou composés non cristallisés qui sont obtenus par voie chimique ou condensation de vapeurs sur support froid.

Il apparaît ici une différenciation entre verres et corps amorphes qui, au fur et à mesure que les études de structures progressent, semble de moins en moins justifiée ou pour le moins inadéquate. On peut ainsi montrer qu'un composé non cristallisé possède la même structure désordonnée (à des petits défauts près) quel que soit le procédé de préparation.

Les nouveaux venus parmi les corps amorphes sont des alliages métalliques

contenant jusqu'à 85 % d'un métal de transition tel que le nickel, le cobalt, le palladium, et un métalloïde tel que le silicium, bore, phosphore. Cette teneur très forte en métal fait que l'on parle souvent à leur sujet de "métaux amorphes". De tels matériaux ont pu être fabriqués grâce à une rapidité très accrue des procédés de condensation de la matière. Par exemple, si de nombreux verres covalents s'obtiennent par un simple refroidissement à la température ambiante du bain fondu, il faut une vitesse de refroidissement évaluée à 10^6 degré/seconde pour empêcher la cristallisation d'un alliage fondu tel que $\text{Pd}^{81} - \text{Si}^{17}$.

Ordre à courte distance

Des travaux récents ont été consacrés à la recherche de l'ordre local dans les "métaux amorphes". Ils aboutissent actuellement à un modèle satisfaisant.

Il existe deux types d'ordre à courte distance : micro-cristaux et désordre homogène.

Laboratoire de Physique des Solides, Orsay (Paris Sud)

propos de toute nouvelle recherche de structure amorphe réside dans le choix du modèle de représentation d'une telle structure. Deux modèles différents sont a priori envisageables : le modèle dit "microcristallin" et le modèle homogène ou "random network" (fig. 1).

• Dans le 1^{er} cas (fig. 1a), le solide est supposé être constitué de cristaux très petits (15 à 20 Å de diamètre). A l'intérieur de chaque microcristal les atomes occupent des positions bien déterminées, identiques à celles qu'ils occuperaient dans un cristal de grande dimension. Par contre, il y a désordre total entre les atomes de cristaux voisins (désorientation des cristaux les uns par rapport aux autres).

• Dans le 2^e cas (fig. 1b), les positions d'un atome et de ses voisins ne sont pas bien déterminées et le désordre s'accroît en passant des premiers aux deuxième, aux troisième voisins, etc... Mais dans ce cas, le matériau est homogène statistiquement, sans solution de continuité ce qui entraîne que les positions de deux atomes très éloignés dépendent des positions des atomes intermédiaires.



1a - poudre de microcristaux



Fig. 1

1b - désordre homogène
"random network"

Schémas à deux dimensions d'un ensemble de 51 atomes de métal regroupés selon deux types de structures différents.

Du point de vue physique, il y a là une différence fondamentale puisque les propriétés électroniques et magnétiques par exemple, dépendent presque exclusivement de la disposition des atomes dans les deux premières couches qui entourent un atome donné.

L'étude des figures de diffraction des rayons X par des échantillons amorphes offre le moyen de choisir entre ces deux modèles à condition d'effectuer des mesures d'intensité très précises et de les exploiter convenablement.

Alliages métalliques amorphes

Les méthodes d'exploitation des diagrammes de corps amorphes ont peu évolué depuis que Debye, puis Zernike et Prins en ont jeté les bases (1915).

Dans un premier temps, on effectue la transformée de Fourier du diagramme expérimental. Le résultat représente la fonction de distribution des atomes dans l'espace, définie comme la probabilité de trouver un nombre N d'atomes à une distance donnée R d'un atome pris comme origine. La structure n'est pas totalement déterminée à ce moment là car on n'a pas d'information sur la manière dont sont disposés les N atomes dans la petite coquille sphérique de largeur Δr autour de R .

On utilise alors les informations recueillies pour bâtir a priori un modèle de structure. Connaissant cette fois la position exacte de tous les atomes, la formule de Debye permet de calculer la fonction d'interférence théorique de ce modèle, qui est comparée au diagramme expérimental. Il faut en général modifier plusieurs fois le modèle pour obtenir un bon accord entre les deux fonctions.

Les diagrammes de diffraction des "métaux amorphes" diffèrent des diagrammes de ces mêmes métaux à l'état liquide. Outre une succession monotone d'anneaux d'intensité décroissante comme en présentent les liquides, il apparaît très souvent un dédoublement sur le deuxième anneau.

En particulier, les alliages nickel-phosphore, palladium-silicium, fer-phosphore-carbone, possèdent exactement le même diagramme (fig. 2). Cela signifie qu'il existe dans ces deux alliages non cristallisés un ordre local identique.

En appliquant la méthode d'investigation décrite précédemment, il s'est avéré que le modèle microcristallin était inapplicable dans ce cas. Il a donc été impossible de trouver un modèle cristallin dont la fonction d'interférence $I(k)$ cal-

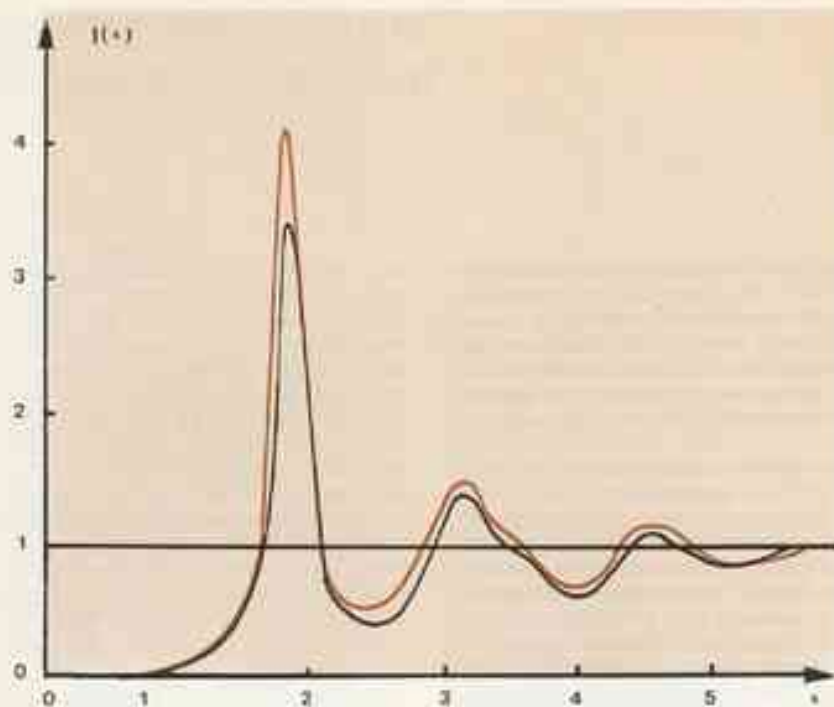


Fig. 2

Fonctions d'interférences comparées des alliages nickel-phosphore amorphe (en couleur) et palladium-silicium amorphe (en noir). Ces deux alliages possèdent le même diagramme.

culée soit identique à la courbe expérimentale. Ce résultat a été confirmé par une étude des quantités de chaleur mises en jeu lors du recuit de ces matériaux (analyse thermique différentielle). À une température de l'ordre de 300 °C, l'alliage amorphe recristallise brusquement en libérant une chaleur presque équivalente à sa propre chaleur latente de fusion. Ceci constitue la preuve que la transformation débute par une germination de petits cristaux, phénomène analogue à ce qui se produit lors du passage liquide-solide. En effet une simple croissance de cristaux ne donnerait pas lieu à un phénomène thermique aussi intense.

Au terme de cette première étape on avait acquis la certitude qu'il pouvait exister de véritables verres à prédominance métallique.

Choix d'un modèle de structure

Un modèle désordonné devait s'appuyer sur les résultats obtenus suivants :

- La densité mesurée du matériau est voisine de la densité d'une structure compacte, cubique, face centrée ou hexagonale compacte.
- L'affinité chimique entre les deux constituants métal et métalloïde est

telle qu'un atome de métalloïde s'entoure préférentiellement d'atomes métalliques.

- Les liaisons entre ces deux éléments ne sont toutefois pas directionnelles comme dans le cas d'une liaison covalente.

Le calcul sur ordinateur permet d'envisager la construction de modèles désordonnés respectant ces hypothèses.

Le programme procède par adjonction à un petit noyau central de nouvelles boules de telle manière que la compacité de l'amas soit maximum. Les boules sont de deux tailles pour simuler les deux types d'atomes. Un sous-programme de génération de nombre au hasard règle l'arrivée sur l'amas des différentes boules. Cependant, chaque fois qu'une petite boule est choisie, le programme fait en sorte qu'elle soit entourée uniquement de grosses boules, ceci pour simuler l'affinité chimique entre éléments métal-non métal.

De nombreux amas contenant chacun un total de mille boules ont été construits ainsi en faisant varier pour chacun, soit la concentration des deux types de boules, soit leur différence de taille. La fonction d'interférence théorique de chaque modèle est calculée puis comparée au diagramme expérimental.

On a superposé sur la fig. 3 la fonction d'interférence de l'alliage nickel-phosphore amorphe et celle calculée pour un modèle contenant 15 % de petites boules de taille inférieure de 10 % à celle des grosses (voir photographie d'une partie du modèle). L'épaule sur le deuxième anneau apparaît pour la première fois dans un calcul de modèle compact désordonné. Le désaccord subsistant entre les intensités relatives du premier anneau sur les deux fonctions provient du nombre insuffisant de boules constituant le modèle. On peut montrer en effet que l'intensité de cet anneau croît très vite avec les premières centaines de boules ajoutées, mais que cet accroissement diminue ensuite sans atteindre cependant une valeur asymptotique pour 1000 boules. Par extrapolation on peut estimer que la valeur expérimentale de l'intensité du 1^{er} anneau serait atteinte pour un nombre de boules au moins égal à 10.000.

L'étude de l'influence de la taille et du pourcentage des petites boules sur la fonction d'interférence calculée montre la sensibilité de l'épaule à ces deux facteurs.

Cette aspérité du deuxième anneau s'atténue si les boules sont de taille par trop différentes (> 15 %) ou bien si la concentration des petites boules se situe en dehors de limites qui vont de 5 à 25 %. Ce résultat est à rapprocher de la constatation expérimentale selon laquelle les phases métalliques amorphes sont obtenues dans des domaines de concentration centrés autour de 15 % en élément non métallique. La structure des alliages métalliques amorphes peut donc être représentée par un empilement d'atomes, désordonné et compact.

Cette structure est difficile à réaliser

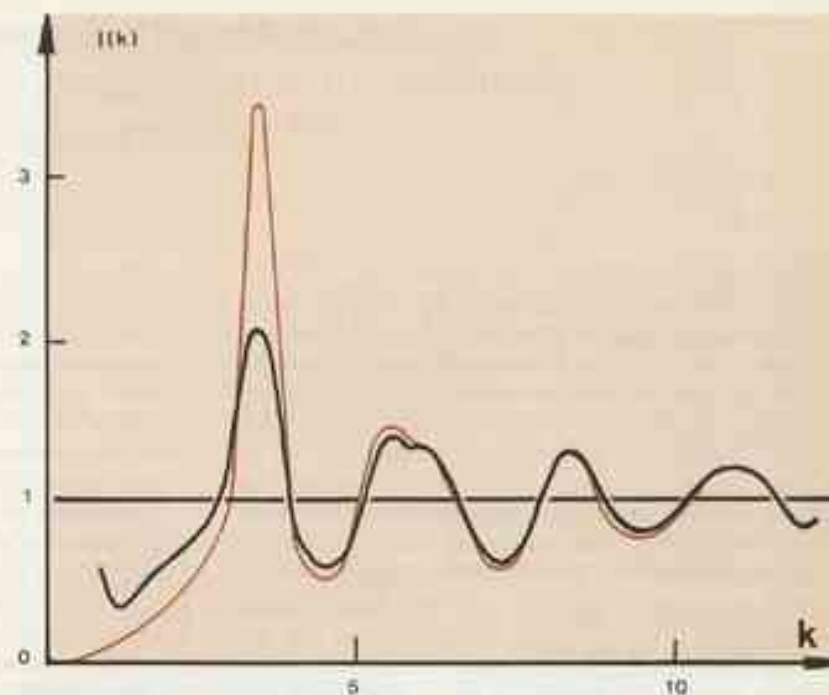


fig. 3

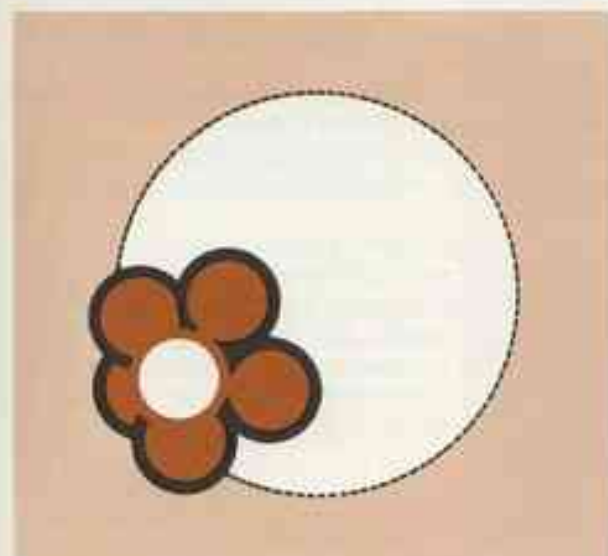
Comparaison du diagramme de la fonction d'interférence théorique de l'alliage nickel-phosphore amorphe (en couleurs) avec le diagramme de la fonction d'interférence calculée (en noir) pour un modèle contenant 15 % de petites boules dont la taille est inférieure de 10 % à celle des grosses.

avec un métal pur car la faible énergie associée à la liaison métallique facilite la diffusion des atomes à l'état solide donc l'apparition de germes cristallins. Par contre, si on disperse 10 à 15 % d'un élément non métallique dans le matériau, celui-ci joue le rôle de fixateur des atomes métalliques et la structure désordonnée est stabilisée.

Les propriétés physiques des "métaux amorphes" sont en cours d'étude. Le magnétisme et les propriétés de trans-

port en rapport avec l'état désordonné font l'objet des principaux développements.

Une application est à signaler pour l'instant : la fabrication de résistances à coefficient de température nul utilisant le fait que la pente de la résistivité peut passer graduellement d'une valeur positive à une valeur négative avec la teneur croissante en élément non métallique.



Exemple d'arrangement existant dans le modèle calculé pour Ni-P. On aperçoit dans le coin gauche un anneau de cinq boules jointes sur le dessus, caractéristique de ce type de structure désordonnée.

les cristaux liquides

Entre l'état cristallin et l'état liquide, on peut rencontrer toute une série de phases de symétries intermédiaires. Ces phases "mésomorphes" ou "cristaux liquides" sont essentiellement obtenues avec des molécules organiques allongées. Les trois phases principales sont caractérisées par l'ordre existant sur la position et la direction des molécules :

- dans les "nématiques", le milieu ordonné parfait est constitué de molécules disposées parallèlement les unes aux autres sans ordre sur leur position ;
- dans les "cholestériques", comme dans les nématiques, les centres de gravité sont répartis au hasard. L'ordre sur l'orientation est caractérisé par une direction privilégiée : l'axe du cholestérique. Les molécules sont toutes parallèles entre elles dans des plans perpendiculaires à l'axe, leur direction tourne quand on se déplace parallèlement à l'axe ;
- dans les "smectiques", il existe, outre un ordre sur l'orientation des molécules, un ordre partiel sur leur position. Les molécules sont réparties en couches : il existe un ordre sur les centres de gravité perpendiculairement à la couche mais pas d'ordre à l'intérieur de la couche.

En pratique, un cristal liquide ne présente pas un ordre aussi parfait que celui que nous avons décrit : comme dans un cristal ordinaire, il existe des défauts. Certains d'entre eux, les dislocations de rotation, ou disclinaisons, ont été observés dès le début du siècle. L'étude de ces défauts, présentée dans un des articles qui suit, est assez complexe. Elle est néanmoins fondamentale pour la compréhension des résultats expérimentaux.

Un autre problème d'importance est la stabilité des matériaux utilisés. Jusqu'à présent, les composés utilisés étaient des "bases de Schiff" peu stables. Des progrès ont été réalisés cette année en France concernant la stabilité, la préparation et les caractéristiques électriques des échantillons nématiques :

- 1) Invention et synthèse (au Collège de France) de matériaux nouveaux, les tolans qui sont des substances très stables et qui joueront un rôle important dans l'avenir.
- 2) Invention d'une méthode pour changer le signe de l'anisotropie de conductivité électrique (LEP).
- 3) Problème de l'ancrage aux parois : pour beaucoup de dispositifs, et pour toutes les expériences fondamentales, il faut imposer une direction fixe

aux molécules nématiques sur les lames de verre qui limitent l'échantillon. Cette direction privilégiée était jadis obtenue par polissage (P. Chatelain). On obtient maintenant des résultats beaucoup plus fiables en utilisant des détergents ou des films évaporés sous incidence oblique.

En ce qui concerne l'hydrodynamique des nématiques, parmi les résultats récents, il faut signaler :

1) La mise en évidence d'instabilités thermiques (phénomène de Bénard) avec des seuils 500 fois plus faibles que pour un liquide isotrope.

2) Les études optiques en présence de champs extérieurs variables :

- déformations au-dessus d'un champ seuil ;
- mouvements de parois.

3) Les études, présentées dans l'article qui suit, des fluctuations spontanées par diffusion laser :

- en volume ;
- en surface.

En ce qui concerne les défauts de structure, on a démontré l'instabilité fondamentale de certaines lignes singulières dans les nématiques. Dans les cholestériques, les défauts sont d'une complexité extrême, mais commencent à être identifiés et interprétés.

Les phases smectiques sont beaucoup moins connues que les nématiques ou les cholestériques mais actuellement :

1) La nature de la phase B qui est "presque" un solide a été considérablement éclaircie par des études aux rayons X sur monodomaines.

2) On a démontré qu'une phase "exotique", la phase H introduite par de Vries, n'était pas différente en réalité de la phase B.

3) Les fluctuations d'orientation dans la phase C ont été étudiées par diffusion laser.

4) On commence à comprendre pourquoi le pas des spirales cholestériques croît quand la température diminue.

D'une façon plus générale, les effets prétransitionnels associés aux changements de phase entre mésomorphes commencent à être analysés (cf article "diffusion Rayleigh"). Les deux articles qui suivent permettent de voir quelques aspects de la nature des travaux actuellement réalisés sur les cristaux liquides et les espoirs qu'ils autorisent.

diffusion Rayleigh inélastique et dynamique des petits mouvements dans les cristaux liquides

Sous l'effet des fluctuations thermiques, la direction locale d'orientation (le directeur) d'un cristal liquide oscille autour de sa position d'équilibre, imposée généralement par les conditions aux limites ou des champs extérieurs. Ces fluctuations d'orientation produisent des basculements du tenseur diélectrique et induisent une très forte diffusion de lumière. Des fluctuations analogues existent dans les fluides ordinaires composés de molécules anisotropes et se voient en diffusion de lumière (les "ailes" de la raie Rayleigh) ; mais ici l'ordre d'orientation à grande distance donne des sections de diffusion (et des temps caractéristiques) de 6 à 8 ordres de grandeur supérieurs.

Diffusion Rayleigh quasi élastique dans les nématiques

Dans un nématique, il existe trois types de déformations : en éventail, de torsion et de flexion (fig. 1), caractérisées par différentes constantes élastiques K_1 , K_2 , K_3 . L'intensité de la lumière diffusée par une fluctuation d'orientation de vecteur d'onde q est de la forme

$\propto Kq^2$ où Kq^2 est le couple de rappel élastique associé à la déformation considérée. La mesure de l'intensité totale de la lumière diffusée permet de déduire la constante élastique K . Mais le spectre des fluctuations du directeur est en principe plus riche d'informations car il renseigne sur la dynamique des petits mouvements dans le cristal liquide. Comment se présente-t-il ? Supposons une déformation angulaire du directeur qu'on laisse évoluer dans le temps. Sous l'effet du couple de rappel élastique, les axes moléculaires retournent à l'équilibre (orientation uniforme) en excitant des tourbillons dans le fluide qui pourraient à leur tour recréer une déformation angulaire du directeur.

Une analyse plus fine montre l'existence de deux temps τ de corrélation différents. Le premier correspond à un mode lent $\tau_1 \approx 10^{-3}$ s (pour $q = 10^3 \text{ cm}^{-1}$), et le second à un mode rapide $\tau_2 \approx 10^{-7}$ s. Malheureusement, ces deux modes ne contribuent pas également à la diffusion de la lumière : l'amplitude des fluctuations est beaucoup plus grande pour le mode lent. La dynamique des fluctuations est alors pratiquement déterminée par un seul temps d'amortissement τ et le spectre est formé d'une seule raie lorentzienne.

L'analyse des spectres de la diffusion Rayleigh, obtenus pour différentes géométries, permet d'accéder aux constantes de viscosité (qui sont au nombre de cinq) du cristal liquide. En effet, la viscosité η qui intervient dans l'expression de la largeur de raie ($1/\tau \sim \frac{Kq^2}{\eta}$) dépend du mode excité (flexion, torsion ou éventail).

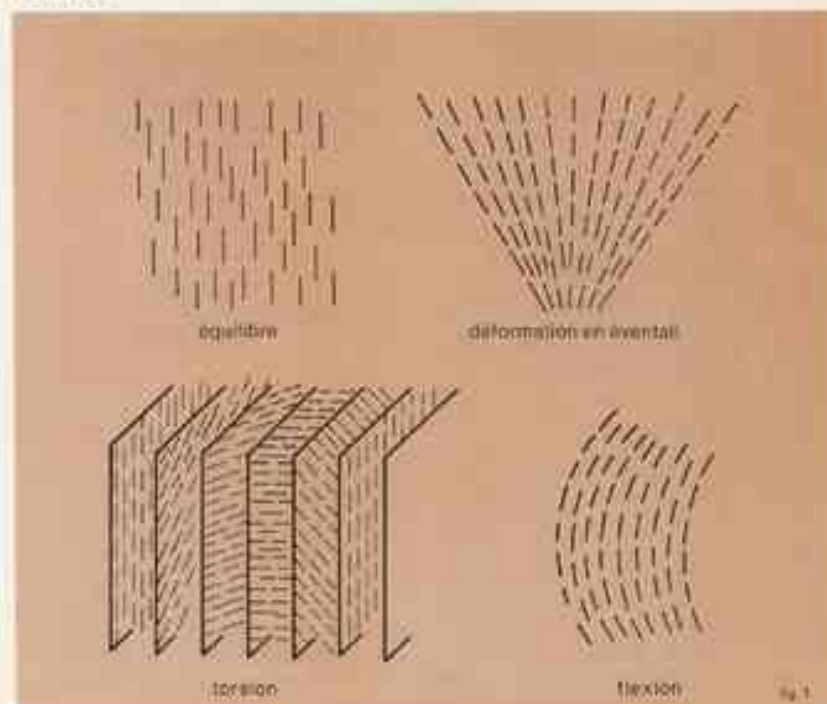
Pour réaliser ces expériences, on utilise des échantillons "monocristallins" orientés par une sorte d'épitaixie liquide : les surfaces limites (verre) sont frottées uniformément ou traitées par divers agents mouillants. Le choix des polarisations de la lumière incidente (ou faisceau laser He-Ne) et de la lumière diffusée permet de séparer les différents modes d'oscillation du directeur. On mesure le temps de corrélation τ par la technique maintenant classique de "battements de photons". Les fluctuations d'intensité de la lumière diffusée (donc du courant photomultiplicateur qui les détecte) sont analysées par un corrélateur en temps réel, qui trace directement la fonction d'autocorrélation $\exp(-t/\tau)$. La dépendance angulaire

Laboratoire de Physique des Solides, Orsay (Paris Sud).

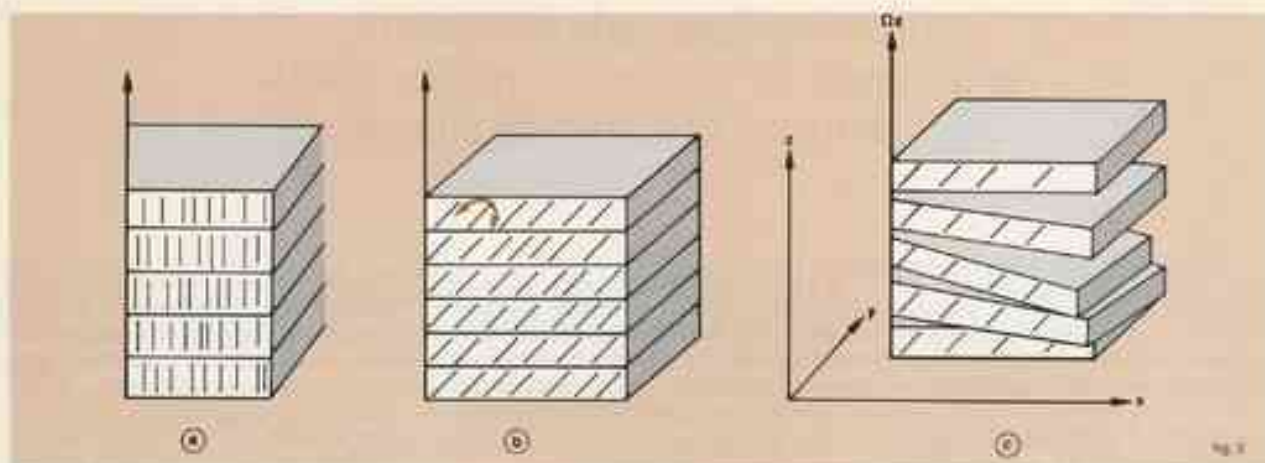
laire de τ a permis d'atteindre jusqu'à présent 4 des 5 viscosités η . On a pu mesurer les constantes élastiques, en superposant au couple de rappel élastique ($\sim Kq^2$) celui produit par un champ extérieur magnétique ou électrique. Ces mesures de τ ont une importance pratique certaine, puisqu'elles permettent de prévoir les temps de réponse des dispositifs électrooptiques utilisant des nématiques. Elles sont aussi plus faciles que les mesures d'intensité.

Diffusion Rayleigh dans les smectiques

Un des problèmes posés par l'élasticité des smectiques est le suivant : les couches peuvent glisser les unes sur les autres sans provoquer de force de rappel élastique. Mais peuvent-elles tourner les unes par rapport aux autres ? Pour les smectiques C, on prévoit que ce type de déformation doit entraîner des couples de rappel élastiques analogues à ceux observés dans les nématiques. L'expérience, réalisée en utilisant la même technique de battements de photons, a mis en évidence l'existence de modes



Représentation : d'un nématique à l'équilibre et des trois types de déformations pouvant y apparaître.



Smectiques A et C

En plus d'un ordre sur l'orientation des molécules analogue à celui d'un nématique, il existe dans un smectique un ordre sur la position de leur centre de gravité. Les molécules sont réparties en couches.

a) - Dans un smectique A, les molécules sont perpendiculaires aux couches.

b) - Dans un smectique C, les molécules sont inclinées par rapport aux couches. L'apparition d'un ordre supplémentaire par rapport à celui existant dans les nématiques (présence de plans smectiques) entraîne la disparition d'un des modes nématiques : celui dont les oscillations (représentées en couleur) changent l'épaisseur des couches.

c) - Mode de torsion plan sur plan existant dans un smectique C.

de torsion plan sur plan (fig. 2), excités thermiquement. Certaines constantes élastiques et de viscosité ont été déterminées pour la première fois dans un smectique C. En comparant sur le même corps, les intensités diffusées en phase C et en phase nématique, on voit de façon spectaculaire l'influence de l'ordre local (apparition des plans smectiques) par la disparition d'un des deux modes nématiques : celui dont les oscillations changent l'épaisseur des couches, qui coûte trop d'énergie élastique en phase C (fig. 2).

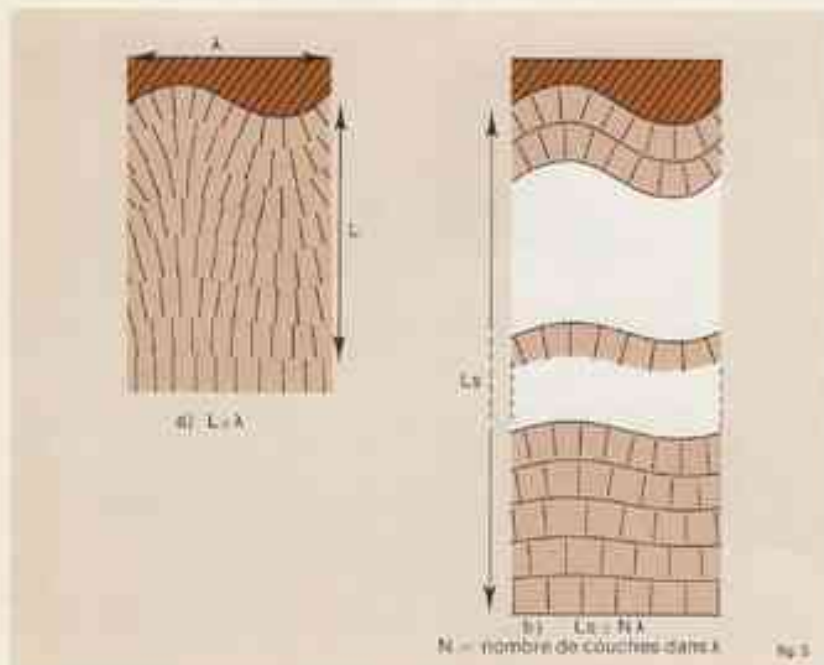
En phase A, on prévoit l'existence d'ondulations thermiques intenses des plans, de vecteur d'onde parallèle aux couches. En fait, l'expérience montre que les couches smectiques reproduisent fidèlement les petites déformations des lames sur lesquelles elles sont déposées sur une épaisseur beaucoup plus grande que la période spatiale de ces déformations. En effet, une déformation périodique (de longueur d'onde λ) de la lame se ressent sur une épaisseur L à l'intérieur du cristal liquide.

Dans un smectique, cette épaisseur est beaucoup plus grande que dans un nématique et s'écrit $L = N\lambda$ où N est le nombre de couches smectiques existant dans λ (fig. 3). Pour un nématique cette profondeur de pénétration est de l'ordre de la longueur d'onde de la déformation. On a là un exemple intéressant de la balance qui existe entre les propriétés cristallines des smectiques, perpendiculairement aux couches, avec leurs propriétés de type cristal liquide dans les couches.

Les transitions de phase dans les smectiques

Entre la phase isotrope et la phase nématique, on voyait déjà des effets pré-transitionnels observés en diffusion de lumière par le groupe du M.I.T. On a montré que certaines transitions de phase $C \rightarrow A$ ou $A \rightarrow$ nématique, qui peuvent être du 2^e ordre, sont isomorphes de certaines transitions de phases plus familières aux physiciens du solide (transition de l'hélium ou supraconducteur $\rightarrow$ normal) mais encore un peu mystérieuses. Il pourrait être possible de mesurer plus facilement certains exposants critiques sur les transitions des smectiques pour les comparer aux calculs théoriques récents (Wilson). Dans cet esprit, plusieurs groupes entament des expériences de diffusion de lumière. A Orsay, on étudie actuellement l'intensité de la lumière diffusée par un nématique à l'approche d'une phase smectique A et on a déjà des indications assez sûres de l'existence d'effets pré-transitionnels.

Si ces efforts aboutissent, on pourra dire que, après être devenues quantitatives, les expériences de physique sur les cristaux liquides, sont peut-être aujourd'hui dignes d'intérêt pour les physiciens traditionnels du solide.



Une déformation périodique (de longueur d'onde λ) de la lame où est déposé un cristal liquide se ressent sur une épaisseur L (profondeur d'atténuation) à l'intérieur du cristal liquide.

a) - Dans un nématique, L est de l'ordre de λ .

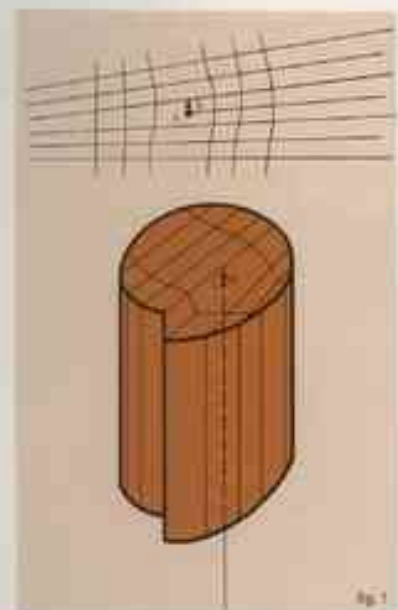
b) - Dans un smectique cette épaisseur est beaucoup plus grande. Elle s'écrit $L = N\lambda$, où N est le nombre de couches smectiques existant dans λ .

dislocations de rotation dans les cristaux liquides

Les défauts de structure d'un milieu ordonné présentent des caractères qui sont représentatifs des éléments de symétrie du milieu. Le cas des dislocations de translation, qui constituent les défauts typiques des solides cristallins, est bien connu : les défauts les plus simples sont alors :

- La dislocation coin (fig. 1a) : la ligne L du défaut constitue la limite d'un plan atomique ; il y a d'un côté de la ligne L un excès (ou un manque) de matière d'épaisseur b correspondant à une translation b , vecteur du réseau, d'un demi-plan par rapport à l'autre.

- la dislocation vis (fig. 1b) : la ligne L du défaut constitue un axe hélicoïdal de pas b , paramètre du réseau, ce qui est visualisé ici pour les plans atomiques perpendiculaires à b .



Dislocation dans un cristal solide.
a) La dislocation coin : à gauche de L il manque une rangée d'atomes, soit une épaisseur de l'ordre de b , paramètre du réseau.
b) La dislocation vis : le cylindre de cristal a été déplacé d'une distance b selon un demi-plan s'arrêtant sur l'axe.

Dans un cas comme dans l'autre on voit que la singularité d'empilement des atomes du cristal est limitée à une petite zone de "mauvais" cristal entourant la ligne, et de dimension de l'ordre de b . Dans toute région extérieure à ce "cœur" les arrangements cristallins, c'est-à-dire les relations entre atomes voisins, sont parfaits, à l'exception de petits déplacements élastiques dus à l'ac-

commodement des contraintes que crée nécessairement la singularité. On sait que ces défauts expliquent les propriétés plastiques des solides cristallins.

Notion de défaut structural dans une phase mésomorphe

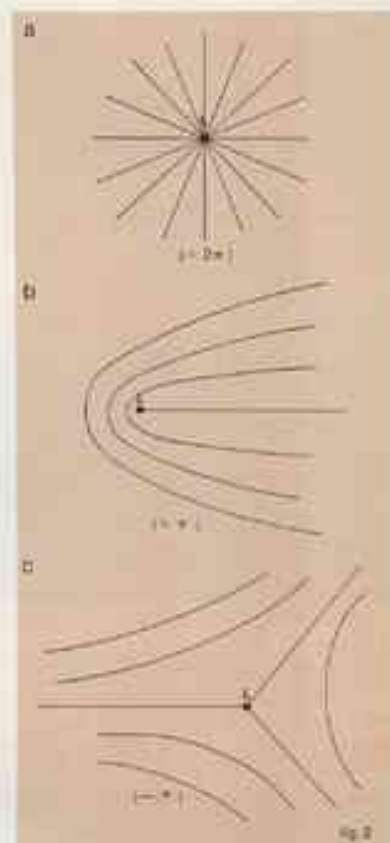
Les phases mésomorphes constituent du point de vue qui nous intéresse non pas des milieux à points (emplacement géométrique des atomes du solide cristallin), mais des milieux à directeurs (direction des molécules). Les défauts qui y apparaissent sont liés aux symétries de rotation (dislocations de rotation encore appelées disclinaisons), et un défaut simple qui y correspond est la dislocation dièdre (fig. 2) : L est rectiligne et parallèle à l'axe de rotation mis en jeu (Frank, 1958). L'ordre $(\pm \frac{1}{2}, \pm 1, \dots)$

d'une dislocation est défini par la rotation que subit le directeur $(\pm \pi, \pm 2\pi, \dots)$ lorsqu'on suit un circuit qui entoure la ligne. On a représenté sur la figure les directions moléculaires, mais on peut encore entendre qu'il s'agit des projections de ces directions. Le terme "dièdre" s'entend facilement : construire une dislocation $+\pi$, par exemple, revient à extraire un coin d'angle π , de sommet L, et à laisser le milieu relaxer de façon à ce que les molécules des faces du coin viennent parfaitement en contact.

Etude des dislocations de rotation

L'étude au microscope polarisant des défauts existant dans les cristaux liquides a permis de déterminer les types d'ordre caractéristiques de ces phases mésomorphes. Les disclinaisons présentes dans un échantillon dépendent des conditions qui existent aux limites de l'échantillon (ancrage des molécules aux parois de l'encainte qui le contient, p.e.), des efforts appliqués sur l'échantillon (champ magnétique), etc. Les disclinaisons ont souvent tendance à s'assembler en textures (Photo 1 en couleur). Ces ensembles joueront nécessairement un rôle dans les propriétés physiques à grande échelle : diffusion de la lumière, effets électriques et magnétiques, viscosité, etc...

• Laboratoire de Physique des Solides, Orsay (Paris-Seul)



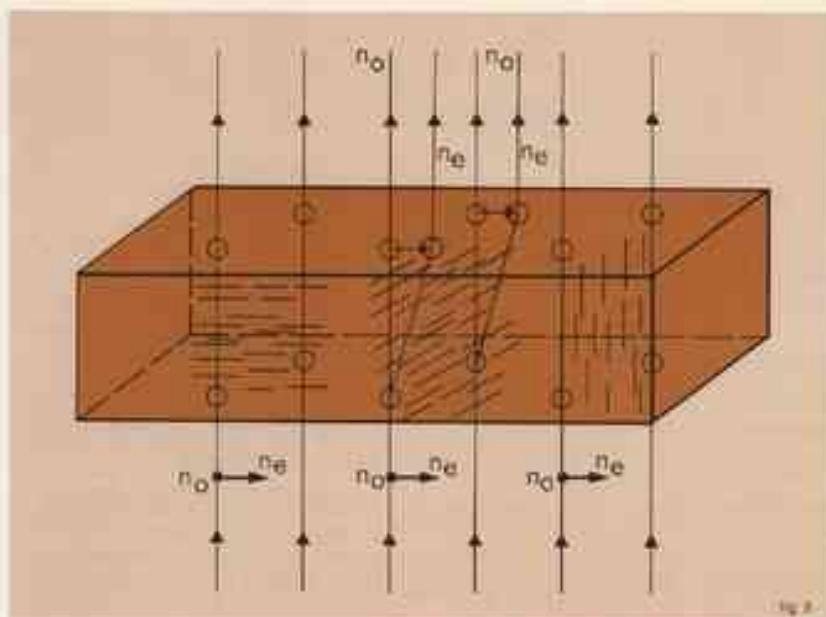
Dislocation dièdre : coupe perpendiculaire à la ligne L. On a représenté les lignes de force des projections des molécules.
a) $S = 1$: le directeur tourne de 2π autour de la ligne ;
b) $S = \frac{1}{2}$: le directeur tourne de π ;
c) $S = -\frac{1}{2}$: le directeur tourne de $-\pi$.

L'étude expérimentale des défauts est du ressort des méthodes d'observation : les cristaux liquides étant biréfringents et transparents, l'instrument adapté est le microscope polarisant. La biréfringence est énorme : $\Delta n \approx 0.3$. La photo 11 en couleur montre des gouttelettes de nématiques germant dans une phase isotrope : on observe des croix de Malte entre nicols croisés, ce qui indique les directions des molécules qui sont parallèles aux nicols, en supposant les molécules parallèles à la préparation. Le nombre des branches noires qui s'échappent d'un "noyau" (selon la terminologie de G. Friedel), ainsi que le nombre de tours effectués par ces branches lorsqu'on tourne la préparation entre nicols fixes, indiquent le rang de cette dislocation.

On utilise aussi des méthodes de décoration qui indiquent directement les déviations moléculaires (Photo III en couleur).

Lorsque les molécules font un angle avec le plan de la préparation, la séparation entre le rayon ordinaire (non dévié) et le rayon extraordinaire (dévié) permet de calculer cet angle et d'avoir une idée sur la direction des molécules : la déviation se fait dans le plan vertical contenant la molécule, le milieu étant uniaxe positif. Si l'on place un réseau de points lumineux sous l'échantillon, le déplacement des images donne une cartographie du directeur (fig. 3), ce qui est illustré sur les photos IV et Va, où une singularité linéaire attachée à la surface sépare deux zones orientées en moyenne à 90° l'une de l'autre. Enfin la polarisation suit quasi adiabatiquement la rotation des molécules autour du vecteur d'onde.

On peut donc mettre en évidence, dans des cas expérimentaux assez simples, les trois types fondamentaux de déformation d'un milieu nématique (flexion, torsion, en éventail). On s'efforce toujours de compléter ces observations des propriétés topologiques des assemblages de directeurs, par des calculs théoriques (minimisation de l'énergie libre de Frank, où interviennent les constantes élastiques K_1, K_2, K_3) de répartition du directeur et d'énergie de ligne. Celle-ci est toujours de l'ordre de K_1 , c'est-à-dire en fait très semblable à celle des lignes de dislocations dans les solides.



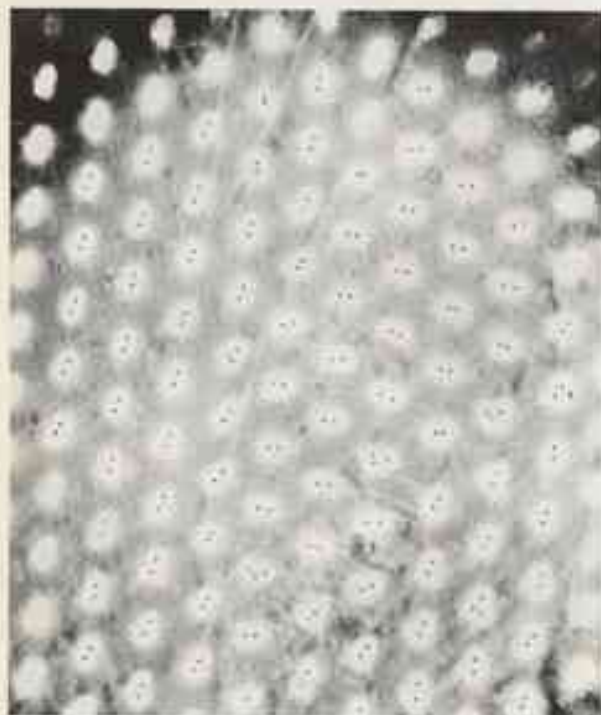
Cartographie du directeur.

On place un réseau de points lumineux sous l'échantillon. Dans le cas où les molécules font un angle différent de 0 ou $\frac{\pi}{2}$ avec le faisceau incident, il y a séparation du trajet ordinaire et du trajet extraordinaire : il s'ensuit un déplacement des images d'un point lumineux.

Lignes de rang unité

Le "cœur" des dislocations dièdres représentées sur la fig. 2 est singulier. Cependant on peut faire subir aux molécules des variations qui, nulle à grande distance de la ligne, contribuent à diminuer au centre l'énergie de cette zone singulière. De ce point de vue, des

arguments topologiques permettent de montrer qu'il est toujours possible de faire disparaître la singularité de cœur d'une ligne de rang entier (multiple de 2π) (fig. 4) alors que c'est impossible pour une ligne de rang demi-entier (multiple impair de π). Ceci explique la fréquence plus grande des lignes de rang unité. On a étudié en détail le cas d'une



a) Réalisation de l'expérience de la fig. 3 avec un réseau de points lumineux.



b) On constate que la déviation des images dépend de la zone choisie, et en particulier que les lignes introduisent des discontinuités dans cette déviation.



I. Structure à noyaux (Schlieren texture) dans une couche mince nématique vue entre polariseurs croisés. (Photo de C. Williams).

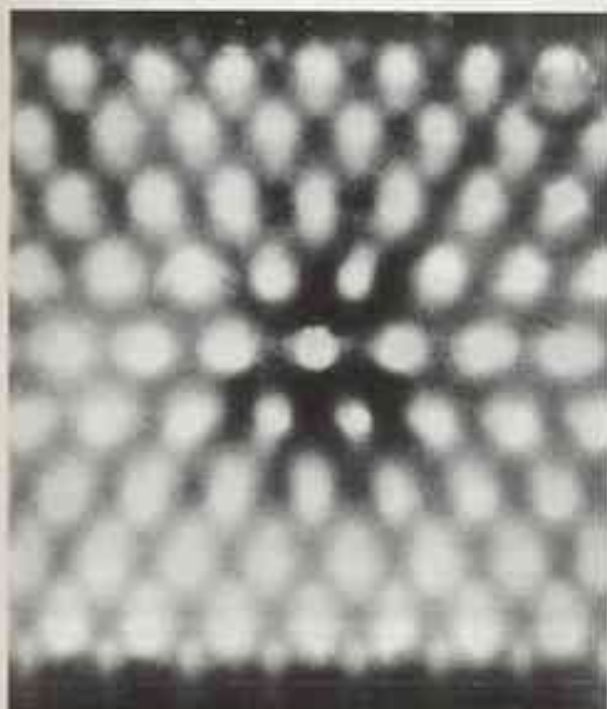
II. Gouttelettes de nématique germant dans la phase isotrope (nucleation) : on observe la présence d'une ligne dièdre au milieu de la gouttelette. (Photo de C. Williams).

III. Décoration des distributions du directeur dans un plan perpendiculaire à une ligne. Les lignes de force sont visualisées par une émulsion de nématique dans un solvant organique. (Photo de P. Cladis et P. Pieranski.)





Fig. 4. Coupe transversale a) et longitudinale b) d'une dislocation dièdre sans singularité de cœur. On a figuré sur la coupe longitudinale deux points singuliers possibles.



Distribution des molécules dans un capillaire à conditions homeotropes. a) Visualisation des directions moléculaires dans un plan méridien par l'utilisation du réseau de points lumineux. On reconnaît l'existence d'un point singulier (cf. fig. 4).

ligne radiale limitée à un cylindre de rayon R , et telle que l'ancrage des molécules est perpendiculaire à la surface extérieure (homeotropie). La fig. 5 représente la variation de l'angle φ en fonction de r ; $A = R \left(\frac{d\varphi}{dr} \right)$ est un paramètre qui décrit les différentes solutions possibles. On a supposé dans le calcul $K_1 = K_2 = K_3$. Les solutions oscillantes de type $A > 1$ (avec retournement du directeur) et de type $A < 1$ sont toujours d'énergie plus forte que la solution limite $A = 1$, où le directeur vient se mettre parallèle à l'axe. Un résultat important est que l'énergie de ligne pour $A = 1$ est $W = \pi K$, c'est-à-dire qu'elle est indépendante du rayon du cylindre.

Après introduction d'un cristal nématique dans un capillaire d'environ 500 μ de diamètre, on a pu observer des lignes de force du directeur par la méthode du réseau de points lumineux (photo Va). La mise au point ayant été faite sur le plan méridien perpendiculaire à l'axe du microscope, on a pu aussi visualiser plus directement ces lignes en utilisant les fluctuations anisotropes de l'indice extraordinaire (photo Vb), qui donnent lieu à la diffusion Rayleigh. Ces fluctuations, observées ici dans l'espace réel, apparaissent comme des taches allongées dans les directions moléculaires. Cette géométrie correspondant à la solution $A = 1$, mais pour d'autres conditions aux limites on peut observer les lignes correspondant aux solutions $A \neq 1$.

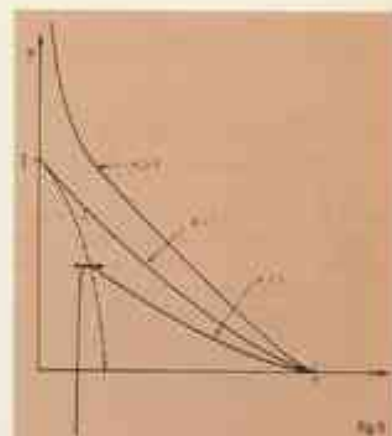
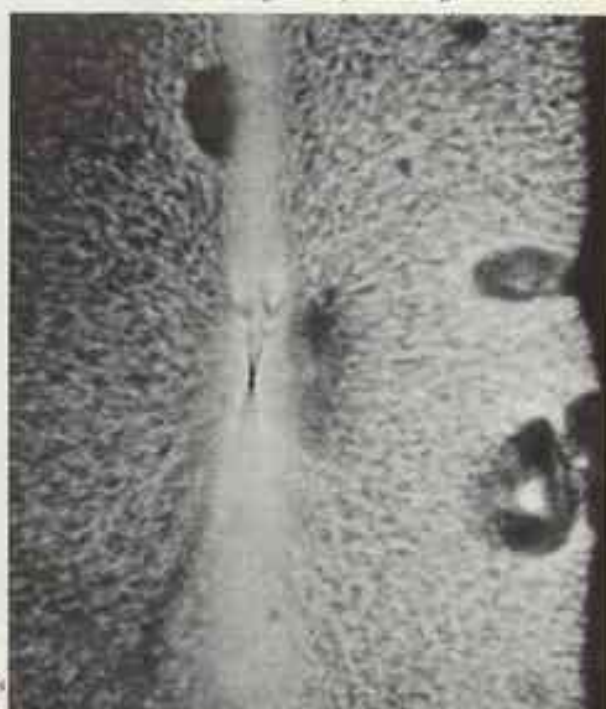


Fig. 5. Variation de l'angle φ du directeur avec un plan perpendiculaire à la ligne, dans la géométrie cylindrique de la fig. 4. La solution $A = 1$ de l'équation de Frank est d'énergie minimale.

Finalement, la courbure des lignes de force peut se faire, soit dans un sens, soit dans l'autre, avec la même énergie. Il doit donc exister des points singuliers sur la ligne, des deux types indiqués sur la fig. 2. On les a étudiés par les mêmes méthodes de visualisation. Ce sont des points singuliers de même nature que ceux qui se présentent dans les gouttes nématiques en phase isotrope, et leur aspect optique est celui des lignes de rang entier. En fait, dans un échantillon de grande dimension les lignes de rang entier n'existent qu'autant qu'elles sont stabilisées par des points singuliers en surface. Entre ces points singuliers, la ligne n'a pas de singularité de cœur.



b) La même région observée à l'aide des fluctuations de l'indice extraordinaire.

Photos C. Williams - P. Pieranski - P. Cladis.

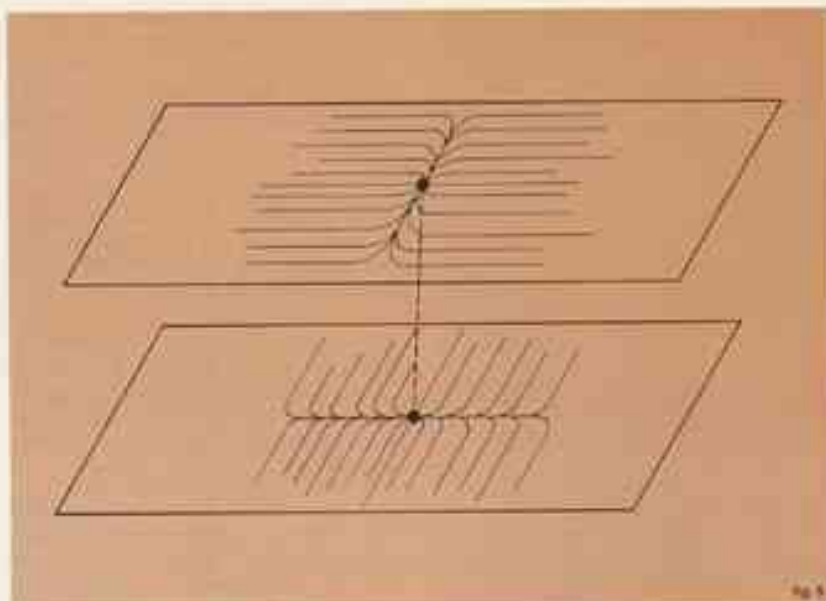
Parois et lignes attachées en surface

Supposons un cristal nématique limité à deux lamelles "frottées", ce qui impose au directeur d'être selon ces directions de frottement. Les molécules aux surfaces, au voisinage de l'émergence d'une dislocation dièdre verticale, ne peuvent pas prendre, par exemple, les directions radiales de la fig. 2. Il y a une déformation de la distribution en quelque chose qui ressemble à une paroi de Néel. La fig. 6 représente cette déformation dans le cas où les deux lamelles sont frottées selon des directions rectangulaires. L'effet observé expérimentalement est illustré sur la photo VI.

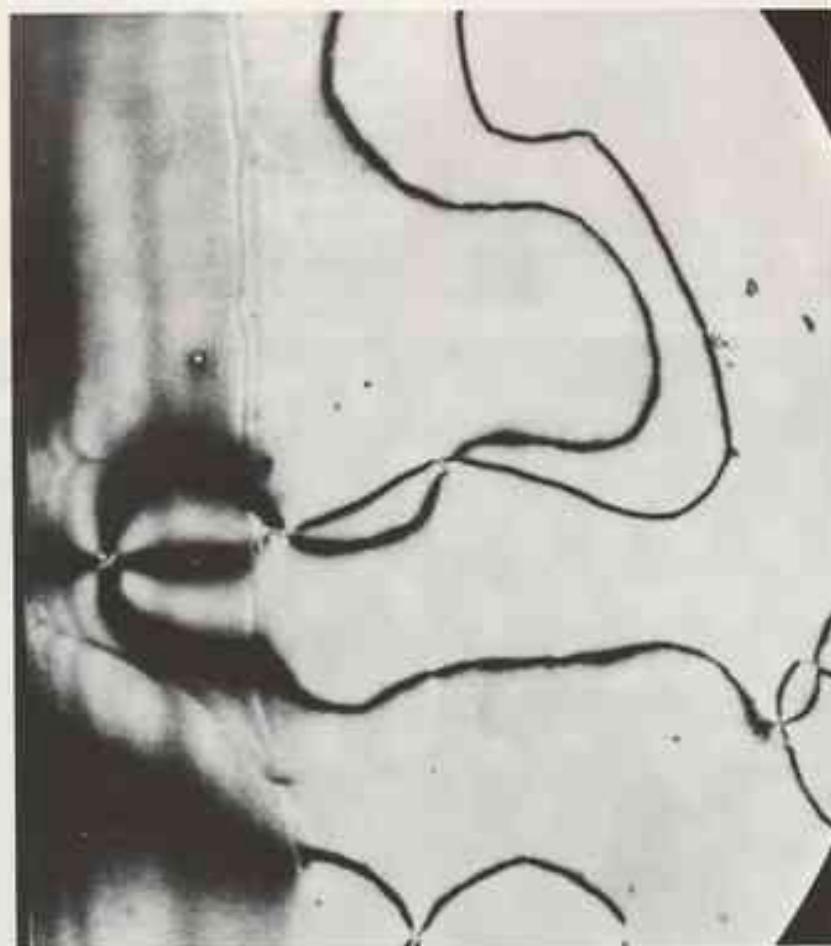
Ces effets de surface doivent effectivement être compris comme des émergences de parois. Dans le cas le plus simple où les lamelles seraient frottées selon des directions parallèles, la configuration se répète sans rotation, et donne une véritable paroi de Néel. Il convient de noter que les dimensions de la paroi mettent en jeu une compétition entre l'énergie de surface et l'énergie de volume. La mesure de la largeur de cette paroi permet donc de mesurer l'énergie de surface.

Cependant il ne serait pas correct d'en déduire que des parois existent dans les nématiques (sauf effets de champ magnétique). Dans le cas d'abord cité de deux frottements rectangulaires, l'expérience montre nettement que la "paroi" qui subit un mouvement hélicoïdal est largement étalée dans toute la masse. Il est sans doute plus juste de dire que les parois se manifestent par leurs émergences en surface, qui sont en réalité des disclinations. On peut s'en assurer, en faisant un circuit autour de l'émergence (sur la fig. 6 le directeur tourne de π). D'ailleurs, si on applique un champ magnétique selon la direction de frottement, l'émergence finit par se détacher et donne lieu dans la masse à une ligne fixe, avec cœur singulier.

Finalement, certaines lignes de surface peuvent conduire à réduire l'énergie d'une situation imposée par les ancrages aux limites. Ainsi, plaçons un nématique entre deux lamelles parallèles, l'une frottée, l'autre telle que le directeur fasse un angle avec la lame, dans un plan ne contenant pas la première direction de frottement. Cette situation de haute énergie peut être relâchée par des lignes de surface. On peut en comparer le rôle à celui joué par des dislocations d'épitaxie, qui relâchent les contraintes qui apparaissent au niveau



Distribution des molécules sur deux lamelles frottées perpendiculairement l'une à l'autre, au voisinage de l'émergence d'une ligne dièdre $S = L$.



Zone nématique entre deux lamelles frottées perpendiculairement l'une à l'autre, s'avancant en coin dans une zone isotrope. On observe la "transition" entre croix de Malte (en coin) et lignes de surface du type de la fig. 6.

photo VI-M. Kleman et C. Williams.

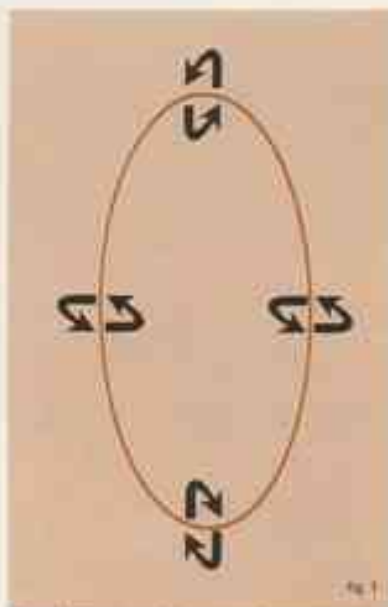
de séparation de deux phases cristallines de paramètres différents; ou bien encore à celui joué par les lignes de vortex dans les supraconducteurs de

deuxième espèce, dont elles diminuent l'énergie par rapport aux supraconducteurs de première espèce, à la même température et dans le même champ.

la vitesse critique dans l'hélium superfluide

A très basse température, l'hélium est un liquide dont plusieurs propriétés sont très différentes de celles des liquides habituels, et ne peuvent s'expliquer que dans le cadre de la mécanique quantique (on dit que l'hélium est un "liquide quantique"). Tout d'abord, de tous les éléments, l'hélium est le seul à demeurer liquide à température absolue T nulle et à pression ordinaire. L'absence de solidification est reliée au phénomène quantique de vibration du point zéro (la relation d'incertitude de Heisenberg interdisant à un atome d'hélium de rester immobile en un point parfaitement défini, il vibre même au zéro absolu, ce qui empêche le solide d'hélium de se former). L'hélium possède deux isotopes stables, de masses atomiques 3 et 4, ^3He et ^4He . L'hélium 4 (de loin le plus abondant des deux isotopes) présente à la température T de 2,18 °K ("point λ ") un phénomène de condensation quantique du type Bose-Einstein. Schématiquement, on peut dire qu'une telle condensation traduit la très forte tendance de toutes les particules à se mettre dans le même état quantique. Ce phénomène a des conséquences importantes : par exemple, le liquide de ^4He peut pour des températures T inférieures à 2,18 °K s'écouler à travers des pores très fins sans "frotter" sur les parois comme le ferait un liquide ordinaire, c'est-à-dire sans gradient de pression ; on dit alors que l'hélium est "superfluide" (Kapitza). La superfluidité tient au fait que, tous les atomes étant dans le même état quantique de vitesse $\vec{v}$ bien définie ($\vec{v}$ peut varier lentement sur des distances macroscopiques), ils sont insensibles aux effets microscopiques qui donnent naissance à la viscosité. La superfluidité est donc, pour des atomes d' ^4He , un peu l'analogue de la supraconductivité pour les électrons groupés en paires dans un métal supraconducteur.

Bien sûr, la température absolue n'est jamais nulle, et il existe toujours dans l'hélium une certaine agitation thermique ; celle-ci se traduit par l'apparition d'excitations dans le liquide (phonons et rotons) qui, elles, peuvent interagir avec la paroi. Pour tenir compte de ces excitations, on définit généralement deux fluides : l'un sans viscosité (le superfluide), l'autre de viscosité normale (le fluide normal), de densités respectives ρ_s et ρ_n (on a constamment



Représentation schématisée d'un tourbillon en anneau. Le "cœur" du tourbillon (anneau représenté en couleur) est constitué de fluide normal ; la circulation, autour de ce cœur, de la vitesse des atomes du superfluide est quantifiée. On peut montrer qu'un tel tourbillon se déplace avec une vitesse perpendiculaire à son plan, inversement proportionnelle à son diamètre.

$\rho_s + \rho_n = \rho$, densité du liquide). Plus le nombre de phonons et de rotons est élevé, plus la proportion du fluide normal dans le superfluide est grande.

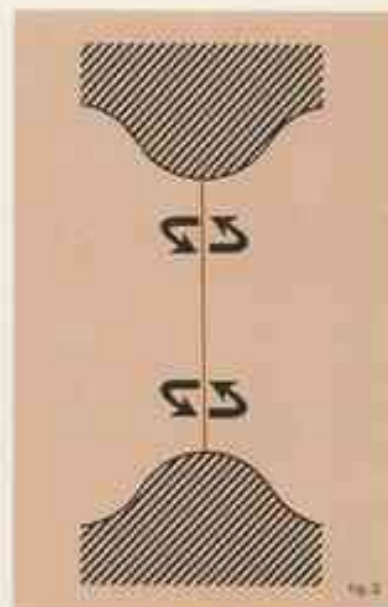
Vitesse critique d'écoulement

Considérons maintenant un superfluide s'écoulant entre deux parois. On s'attend à ce que l'écoulement sans dissipation disparaisse, si la vitesse d'écoulement dépasse une certaine valeur, appelée "vitesse critique". Deux types de transitions quantiques peuvent expliquer l'apparition d'une vitesse critique :

- création d'excitations thermiques, c'est-à-dire diminution du nombre de particules dans l'état superfluide. Le modèle correspondant, donné par Landau, prédit l'existence d'une vitesse critique V_c , due à l'émission de rotons ; le schéma de dissipation est alors :

superfluide $\xrightarrow{V_c}$ rotons $\rightarrow$ parois

Laboratoire de Physique - Ecole Normale Supérieure.



Tourbillon linéaire piégé sur les parois à ses deux extrémités.

Cependant, la vitesse critique correspondante est de l'ordre de 6×10^3 cm/s et, pour le moment, une telle vitesse n'a jamais été observée.

- création de tourbillons en forme d'anneau (Feynman, Onsager), le nombre de particules dans l'état condensé restant cette fois constant (fig. 1). Le cœur d'un tel tourbillon (analogue d'un rond de fumée) est constitué par du fluide normal qui, lorsque le tourbillon se déplace, peut "frotter" sur le liquide normal du bain. Le schéma de dissipation est alors :

superfluide $\xrightarrow{V_c}$ tourbillon $\rightarrow$
frottement sur le fluide normal $\rightarrow$ excitations thermiques
 $\rightarrow$ parois

Il existe d'ailleurs d'autres types de tourbillons, tels que les tourbillons accrochés par leurs deux extrémités à la paroi (figure 2). Cette fois encore, leur cœur est constitué de fluide normal ; lorsqu'ils sont placés dans un courant d'hélium liquide, ils peuvent se déplacer, ce qui donne lieu à des processus de dissipation.

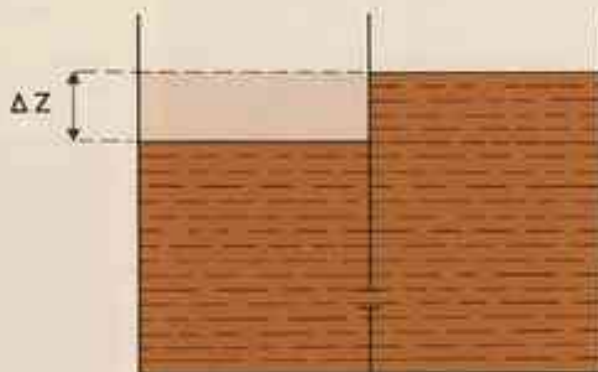


Fig. 3

Schéma de l'expérience. Les deux vases, remplis d'hélium à très basse température, communiquent par un trou très petit, et on étudie le retour à l'équilibre de la différence de niveau ΔZ .

Etude expérimentale de la vitesse critique

Les expériences qui vont suivre concernent l'étude de la vitesse critique de l'écoulement d'hélium liquide à travers un trou de très petite dimension.

Le principe en est très simple (figure 3) : on remplit d'hélium liquide à très basse température ($0,3 \text{ K} \leq T \leq 2,17 \text{ K}$) deux récipients communiquant par un orifice dont le diamètre est compris entre 1 et 20 microns. En enfonçant ou en soulevant un flotteur dans l'un des deux récipients, on impose une différence de niveau ΔZ entre les deux bains couplés, et on étudie ensuite le retour à l'équilibre. Pour un liquide normal, la vitesse d'écoulement serait proportionnelle à la différence de pression des deux côtés de l'orifice, c'est-à-dire à ΔZ , et le retour à l'équilibre donnerait pour ΔZ une variation en exponentielle amortie. Pour l'hélium superfluide, les choses se présentent de façon complètement différente : on observe une vitesse d'écoulement constante (égale à V_c) quelle que soit la différence de hauteur ΔZ (fig. 4). Lorsque l'équilibre est atteint ($\Delta Z \approx 0$), des oscillations apparaissent (le superfluide oscille entre les deux bains) très faiblement amorties (elles peuvent durer plusieurs heures).

La mesure de la pente de la droite qui donne ΔZ en fonction du temps t fournit la vitesse critique. Si le bain d'hélium utilisé conserve la turbulence qui s'établit inévitablement lors du remplissage

des récipients, la valeur mesurée pour V_c est de l'ordre de 40 cm/s ; on interprète cette valeur comme provenant de l'existence d'un grand nombre de tourbillons linéaires piégés dans le trou. Il est toutefois possible d'éviter la formation de tels tourbillons ; il suffit pour cela, lors du remplissage, de placer en face du trou un filtre comprenant un grand nombre de pores de petites dimensions (ce filtre est constitué d'une poudre fine comprimée) ; la turbulence du liquide est de ce fait filtrée. Effectivement, la vitesse critique mesurée est alors beaucoup plus élevée, de l'ordre de 1000 cm/s . Cette vitesse critique peut être reliée à la création de tourbillons en anneau, qui germent dans le fluide sous l'effet des fluctuations thermiques et, lorsqu'ils sont placés dans un courant liquide, tendent à grandir. Pour donner une analogie classique, on peut dire que l'apparition d'une vitesse critique par les deux processus envisagés est semblable à la formation des gouttes dans une vapeur sursaturante, soit sur la paroi, soit en volume par fluctuation de température.

L'étude des oscillations de ΔZ au voisinage de l'équilibre a également été entreprise ; en mesurant l'amplitude maximale de ces oscillations, on peut déterminer la vitesse la plus grande que peut avoir le liquide circulant dans le trou, ce qui permet d'obtenir une valeur indépendante pour V_c . Il a ainsi été possible de confirmer les résultats précédents.

Le prolongement naturel de ces expériences est de travailler à des tempé-

ratures encore plus basses, de façon à réduire les fluctuations thermiques et augmenter ainsi la vitesse critique. On peut même espérer atteindre les vitesses très élevées où ce sont les créations directes d'excitations thermiques qui limiteront V_c .

Stabilisation de la différence de niveau entre les bains

Rappelons que, dans notre liquide quantique, toutes les particules sont dans le même état quantique. Par suite, de même que dans un atome le moment cinétique d'un électron qui tourne autour d'un noyau est quantifié (la quantification de Bohr-Sommerfeld indique que ce moment est un multiple entier de $\hbar$, constante de Planck divisée par 2π), la circulation de l'hélium autour de la discontinuité qui constitue le cœur du tourbillon est quantifiée :

$$\oint \vec{v} \cdot d\vec{l} = \frac{h}{m}$$

(m est la masse d'un atome ^4He ; le contour d'intégration de la vitesse est représenté sur la figure 1). Par analogie avec l'effet Josephson alternatif dans les supraconducteurs, on peut essayer de mettre en évidence les effets quantiques associés (effet Anderson-Richards, du nom des physiciens qui, les premiers, l'ont observé). Soit ν la fréquence de passage des tourbillons dans le trou ; la quantification de la circulation dans chaque tourbillon permet, en écrivant la conservation de l'impulsion, de relier la différence de

pression Δp de part et d'autre du trou à la fréquence ν par :

$$\Delta p = \frac{h}{m} \nu$$

Δp est évidemment proportionnel à ΔZ , et on a l'égalité

$$\Delta Z = \frac{h}{mg} \nu$$

où g est l'accélération de la pesanteur. Par exemple, pour une différence de niveau ΔZ de 1 mm, cette relation donne $\nu = 10^5$ Hz. A la suite de Anderson et Richards, cette émission de tourbillons a pu être synchronisée en plaçant, en face du trou, un quartz vibrant à la fréquence ν_0 ; on montre alors qu'il y

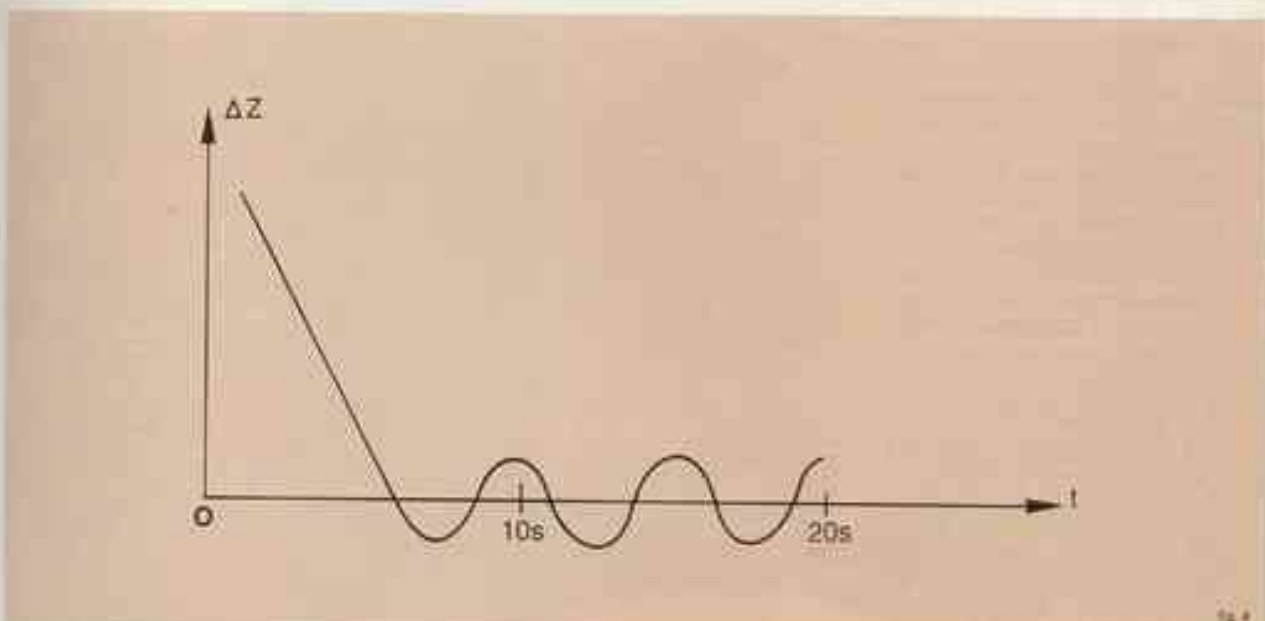
a stabilisation de la différence de hauteur ΔZ chaque fois que $\nu = n \nu_0$, où n est un entier. Effectivement, des courbes telles que celle de la figure 3 ont pu être obtenues : on constate que ΔZ se stabilise autour de valeurs non nulles, multiples de $h \nu_0 / mg$; pour passer d'une valeur stable de ΔZ à la suivante, on perturbe le système.

Il faut cependant signaler que l'analyse d'une telle expérience est très compliquée, et que l'interprétation précédente est actuellement discutée. Seules les expériences qui vont être poursuivies

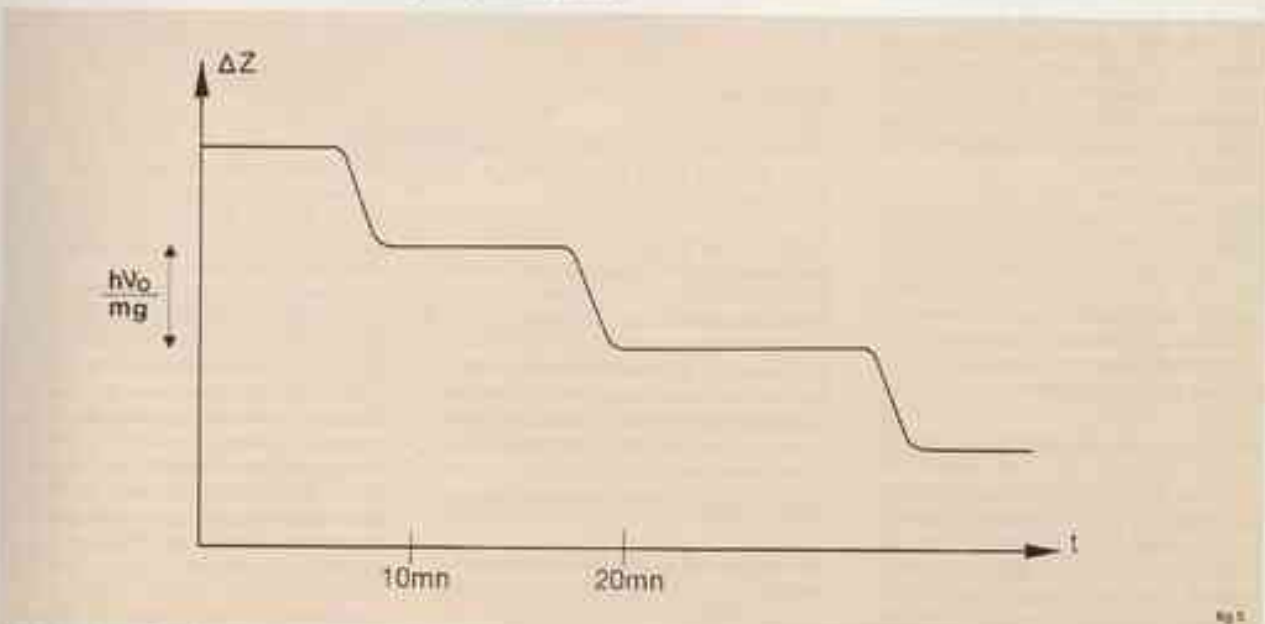
permettront de trancher entre les diverses possibilités.

Aujourd'hui, l'apport le plus essentiel de ces expériences semble être de fournir les premières données dans l'étude des liquides, étude qui malgré son importance est encore dans l'enfance du fait de son extrême difficulté.

Il est trop tôt pour affirmer que ce type de recherche trouvera des applications aussi importantes que celles de l'effet Josephson. Quoi qu'il en soit, l'expérience est très belle et autorise des espoirs d'applications pratiques par exemple pour des mesures de g .



Variations en fonction du temps de la différence de niveau ΔZ . On obtient d'abord une droite, dont la pente donne la vitesse critique. Une fois l'équilibre atteint, on observe des oscillations pendant plusieurs heures.



Variations de ΔZ en fonction du temps t , lorsque un quartz de fréquence ν_0 est positionné en face du trou : la différence de niveau se stabilise à différentes valeurs, et on peut leur faire correspondre la relation $mg \Delta Z = n h \nu_0$.

réfrigérateur à dissolution et physique des basses températures

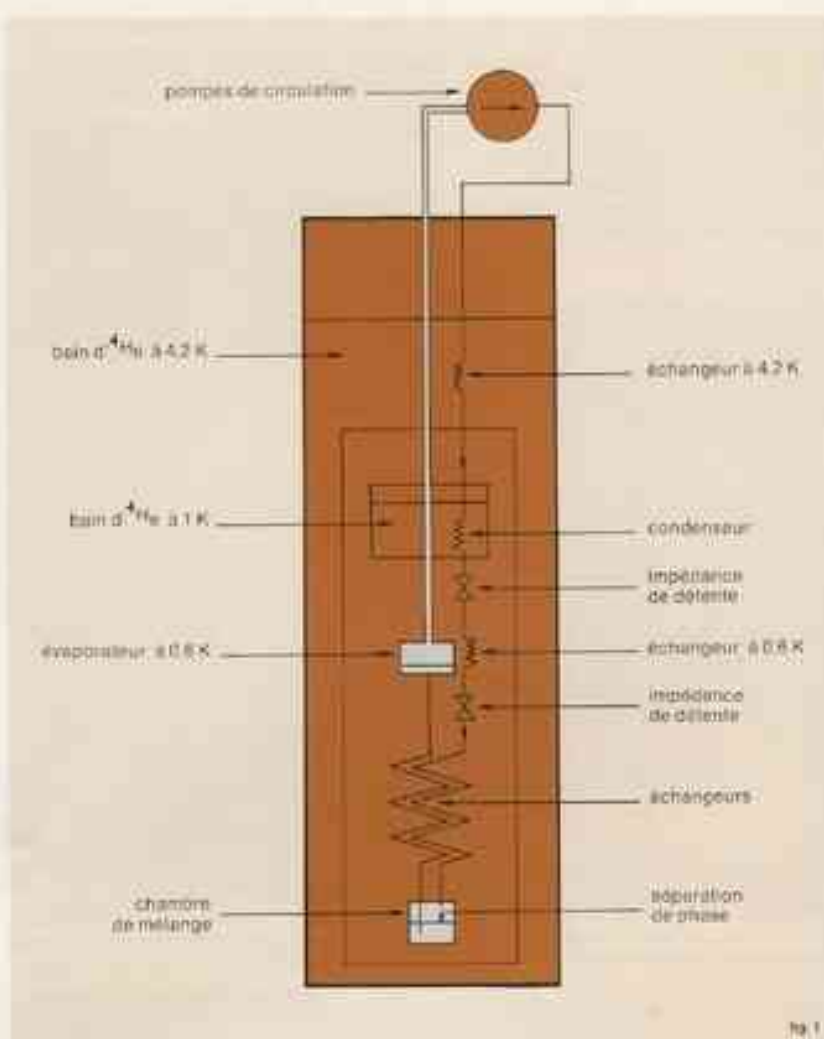
Les recherches en physique expérimentale aux températures inférieures à 1 Kelvin connaissent un essor considérable, motivé notamment par l'emploi de plus en plus répandu de réfrigérateurs à dissolution d'hélium-3 dans l'hélium-4. Ces réfrigérateurs permettent d'atteindre et de maintenir des températures de l'ordre de 10^{-3} Kelvin (10 milliKelvin) aussi longtemps que nécessaire et d'absorber des puissances thermiques qui vont jusqu'à des fractions de milliwatt vers 10^{-3} Kelvin.

Ces réfrigérateurs trouvent leur application dans de nombreux domaines de la physique expérimentale. En physique nucléaire, ils sont utilisés pour refroidir des cibles polarisées, soit par polarisation dynamique, soit par force brutale (par application d'un champ magnétique intense). Cette dernière méthode est grandement facilitée par le fait que les réfrigérateurs à dissolution sont très peu sensibles au champ magnétique. On peut ainsi obtenir des rapports H/T (champ magnétique sur température) pouvant atteindre 1000 Tesla par Kelvin. Les expériences d'effet Mössbauer et les mesures d'anisotropie de rayonnement gamma à très basse température sont également facilitées par le fait que l'on peut disposer d'une puissance frigorifique stable des jours durant. Il faut également noter les études en cours sur les propriétés de l'hélium liquide (isotopes 3 et 4 et mélanges) à des états quantiques très voisins de l'état fondamental. Il en est de même du problème des impuretés magnétiques dans les métaux, ainsi que dans de nombreux autres domaines de la physique de la matière condensée où la détermination des propriétés de l'état fondamental est d'importance.

De nouvelles perspectives s'ouvrent enfin, en cosmologie expérimentale, avec la construction conjointe à Stanford, Baton Rouge et Rome, de détecteurs d'onde de gravité, constitués de résonateurs en aluminium de cinq tonnes refroidis par dissolution au-dessous de 100 milliKelvin.

Principe de fonctionnement

Le principe de fonctionnement des réfrigérateurs à dissolution d'hélium-3 dans l'hélium-4 est fondé sur les propriétés des solutions isotopiques de l'hélium. Les isotopes de masse 3 et 4 de l'hélium restent liquides jusqu'au zéro absolu sous leur tension de vapeur saturante. Cette propriété bien connue,



Réfrigérateur à dissolution continu.

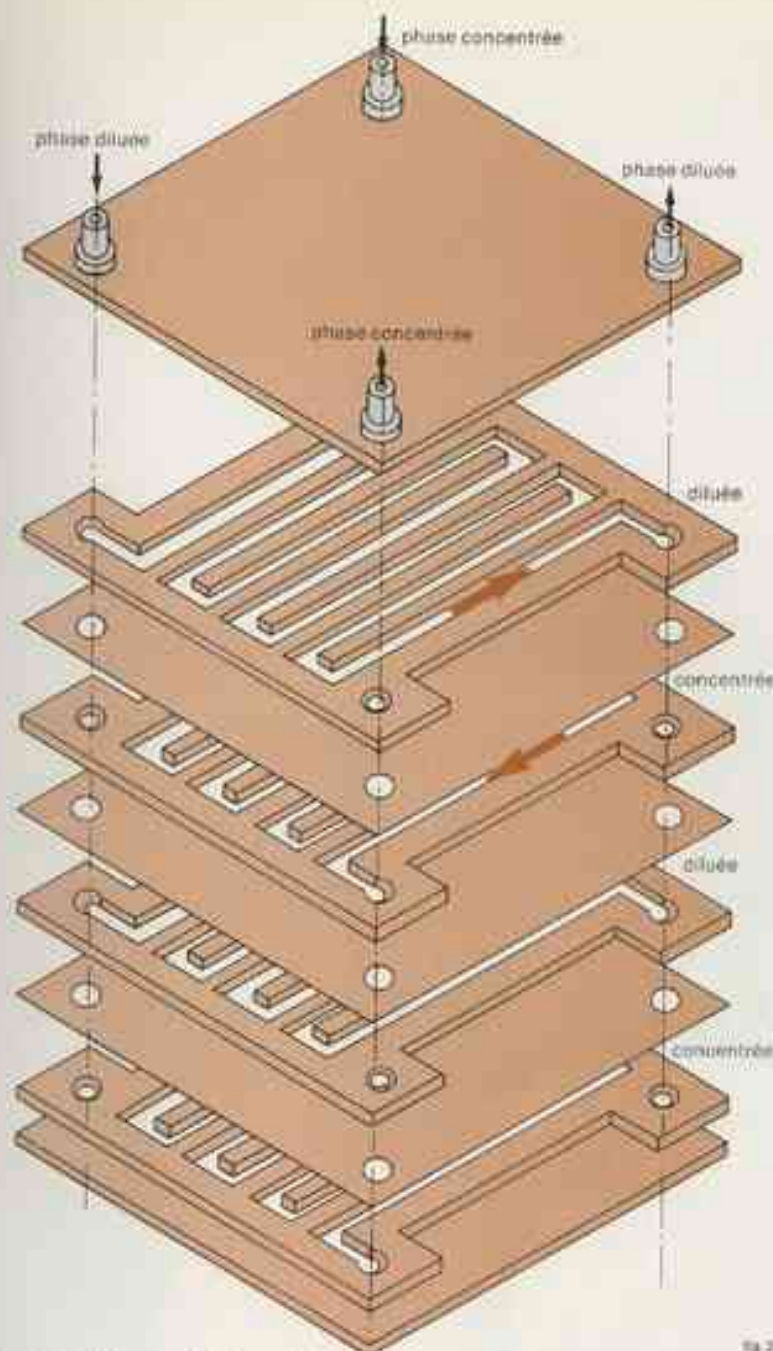
résulte de la faiblesse des interactions entre atomes et de la grande légèreté de leurs noyaux.

Les solutions d'hélium-3 et d'hélium-4 restent elles aussi liquides jusqu'au zéro absolu, mais elles ne sont inconditionnellement stables qu'au-dessus de 0,87 Kelvin. En-dessous de cette température, les mélanges des deux isotopes 3 et 4 de l'hélium se séparent spontanément en deux phases (démixion). L'une, dite phase concentrée, est composée d'hélium-3 presque pur. L'autre, dite phase diluée (dont la concentration molaire

s'abaisse jusqu'à 6,5 % d'hélium-3 au zéro absolu), se comporte dans une certaine mesure dans l'hélium-4 comme un gaz parfait dans le vide. Ce "quasi gaz" d'hélium-3 est en équilibre avec l'hélium-3 presque pur de la phase concentrée, de façon analogue à l'équilibre liquide-vapeur, la tension de vapeur mise en jeu étant ici la pression osmotique de l'hélium-3 dans l'hélium-4 superfluide. Si maintenant, des atomes d'hélium-3 sont enlevés à la phase diluée, d'autres atomes de la phase concentrée traverseront la surface de séparation et seront dilués de façon à rétablir l'équilibre. Ce processus s'accompagne d'une absorption de chaleur et l'ensemble se refroidit.

L'existence d'une chaleur de dissolution importante explique l'efficacité de la production du froid par dissolution

- Centre de Recherche sur les très basses températures - Grenoble.
- Institut d'Electronique Fondamentale, Orsay (Paris Sud).
- Laboratoire de Physique des Solides, Orsay (Paris Sud).



Dans ce type d'échangeur, l'échange s'effectue au travers de feuilles très minces de Kapton. Ces feuilles sont collées sur des espaces de Kapton de l'ordre de 0,3 mm d'épaisseur dans lesquels sont dirigés les vannes en chicane parcourus par le fluide. Les dimensions de ces canaux déterminent l'impédance et la surface d'échange élémentaire (9 cm² par feuille de Kapton dans ce cas). On obtient des échangeurs élémentaires en collant en parallèle quelques-uns de ces circuits dans lesquels chaque phase circule alternativement à contre-courant. De tels éléments peuvent ensuite être reliés en série pour obtenir la surface totale voulue.

de l'hélium-3 dans l'hélium-4 : il est en effet possible d'atteindre la température extrêmement basse de 4 milli-Kelvin. Toutefois la caractéristique la plus intéressante de ce cycle de refroidissement réside dans le fait qu'il peut être mis en œuvre d'une manière continue. Des températures inférieures à 10 milli-Kelvin peuvent être maintenues d'une manière pratiquement indéfinie.

Cependant il existe une limite aux températures minimales obtenues par

cette méthode si bien qu'il ne semble pas que l'on puisse actuellement atteindre des températures très inférieures à 5 milli-Kelvin.

Réfrigérateur à dissolution

On a pu réaliser sur ce principe un réfrigérateur à dissolution continue (fig. 1). Partant de la "chambre de mélange" où s'effectue la dissolution, l'hélium-3 dilué remonte à travers un échangeur

jusqu'à l'évaporateur où il est distillé sous pression réduite, la pompe se trouvant à la température ambiante. La température de l'évaporateur est maintenue au voisinage de 0,6 Kelvin, de sorte que la tension de vapeur de l'hélium-4 soit très faible, mais celle de l'hélium-3 assez élevée pour qu'on puisse l'extraire presque pur et avec un débit suffisant. L'hélium-3 gazeux est ensuite réinjecté du côté concentré, refroidi à 4 Kelvin, condensé à 1 Kelvin, encore pré-refroidi dans les échangeurs de chaleur. Ces échangeurs à contre-courant tirent parti de la chaleur spécifique du "quasi-gaz" qui remonte froid de la chambre de mélange. L'hélium-3 est finalement introduit à nouveau dans la chambre de mélange et le cycle peut continuer avec un bon rendement thermodynamique.

Tous les éléments constitutifs, boîte de mélange, tubes de connexion, évaporateur, ont leur importance pour le bon fonctionnement du réfrigérateur, mais ce sont avant tout les échangeurs de chaleur qui en déterminent les performances ultimes. On montre en effet que la température minimale que l'on peut atteindre dans la boîte à mélange est le tiers de celle du fluide concentré à sa sortie de l'échangeur.

L'efficacité de l'échange de chaleur entre les deux phases, diluée et concentrée, est limitée par l'existence à l'interface solide-hélium d'une forte résistance thermique appelée résistance de Kapitza, qui varie de façon inversement proportionnelle au cube de la température. Pour échanger la puissance nécessaire il faut augmenter la surface d'échange tout en minimisant le volume et les pertes de charge (ces dernières afin de ne pas provoquer de réchauffement du fluide par frottement visqueux).

Pour obtenir de grandes surfaces d'échange et de faibles impédances, on choisit habituellement, à cause de leur bonne conductivité thermique, des métaux purs sous forme de feuilles ou de poudre frittée. Ces échangeurs sont isothermes, et discrets, par opposition aux échangeurs dits continus. C'est la solution qui avait été adoptée à Grenoble, Orsay et Saclay dès 1966. Cependant, il a été montré à Orsay que dans ces métaux finement divisés d'autres résistances thermiques s'ajoutent à la résistance de Kapitza, diminuant d'autant l'échange de chaleur.

Pour une surface donnée, l'échange peut encore être amélioré en utilisant des matériaux présentant une résistance de Kapitza plus faible. Des matières plastiques sont utilisées à Grenoble car leur résistance de Kapitza est environ sept fois moindre que celle du cuivre pur. Malgré la mauvaise conductivité thermique de ce matériau, l'échange entre les deux phases d'hélium reste possible si les feuilles qui les séparent sont suffisamment fines pour que leur résistance thermique soit faible par rapport à la résistance de Kapitza. De plus, cette mauvaise conductivité thermique rend possible la réalisation d'échangeurs du type continu, plus efficaces que les échangeurs discrets. On a réalisé des échangeurs dans lesquels l'échange s'effectue au travers de feuilles très minces de kapton de 7 μ d'épaisseur, dont la résistance thermique est 1 % seulement de la résistance de Kapitza. L'arrangement adopté permet d'obtenir des échangeurs continus de grande surface d'échange, tout en conservant de faibles impédances (fig. 2). L'avantage essentiel de cet arrangement est de pouvoir adapter indépendamment l'impédance et la surface d'échange.

Avec un échangeur de ce type (d'une surface d'échange de 100 cm²), comportant 3 éléments en série, lui-même en série avec un échangeur tubulaire classique (fig. 3), on peut atteindre, en régime continu 12 à 16 millikelvin dans la boîte à mélange, pour des débits de 10 à 80 micromoles/sec. d'hélium-3, ceci malgré des apports de chaleur parasite de 2 ergs/sec. Sans réinjecter l'hélium-3, on atteint 7 mK. Avec des échangeurs classiques en cuivre fritté, ceci ne serait possible qu'avec des surfaces au moins dix fois plus importantes. Les performances obtenues rejoignent celles de l'échangeur continu idéal. Avec une surface raisonnable de 1000 cm², il serait possible d'atteindre environ 6 millikelvin en régime continu, en tenant compte de la même puissance parasite.

Ces échangeurs sont faciles à construire et ne posent pas de problème particulier d'étanchéité mais ils se révèlent quelque peu sensibles aux surpressions. L'utilisation de feuilles un peu plus épaisses permet de s'affranchir de cet inconvénient sans que les performances en soient trop affectées.

Outre leurs applications dans de nombreux domaines de la Physique, les réfrigérateurs à dissolution sont surtout

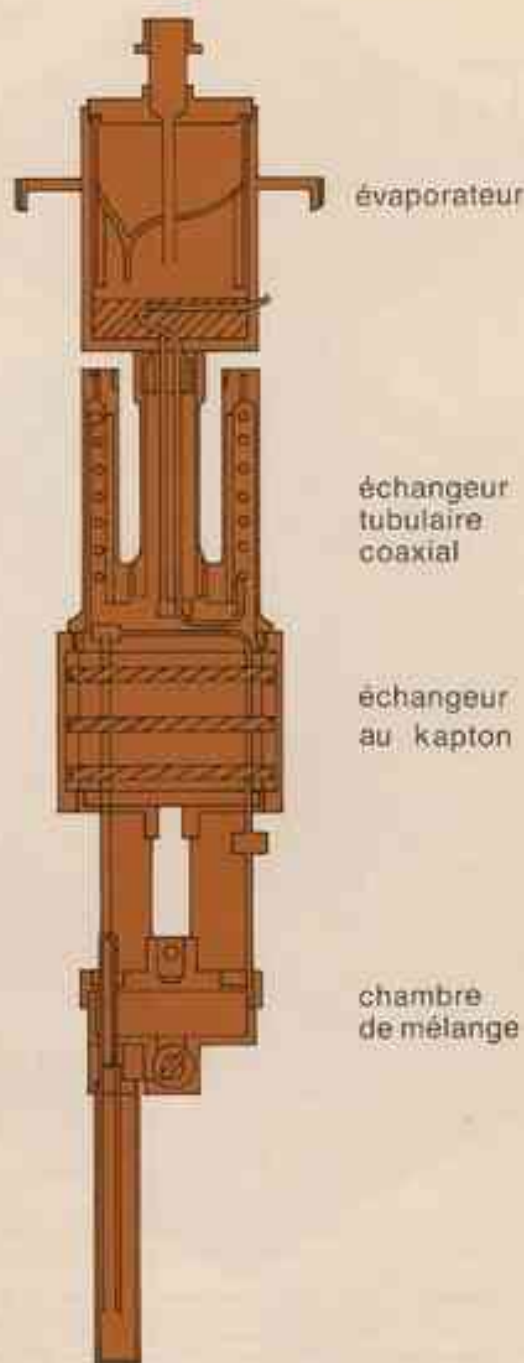


Fig. 3

Cet échangeur comporte trois éléments en série de 18, 26 et 80 cm². Il est lui-même placé en série avec un échangeur tubulaire classique d'une surface équivalente de 3 cm².

employés à l'heure actuelle dans les laboratoires spécialisés en cryogénie. A Saclay sont réalisées de très belles expériences sur l'antiferromagnétisme nucléaire, à Orsay des études sur les propriétés fondamentales de l'hélium liquide et la thermométrie. A Grenoble enfin, des recherches sur les transferts de chaleur, la conduction thermique,

et les propriétés des alliages magnétiques dilués pour la mesure de la résistivité, de l'aimantation sont effectuées. Les réfrigérateurs à dissolution sont également utilisés dans d'autres laboratoires de cryogénie, un peu partout dans le monde, comme point de départ à l'obtention de températures encore plus basses.

sciences pour l'ingénieur



Le circuit électronique (Photo du L.A.S.)

les couches très minces de polymère

L'étude de la formation de couches très minces de polymères et de leurs propriétés électriques fait l'objet depuis quatre ans de recherches au Laboratoire de Génie Électrique de l'université Paul Sabatier de Toulouse.

Ces films sont formés actuellement par polymérisation de vapeurs de styrène dans une décharge électrique alternative. Le dispositif utilisé est le suivant : dans une cloche où règne un vide de 10^{-4} torr, on introduit le styrène à l'aide d'un tube capillaire calibré qui permet d'obtenir une pression de gaz constante. La pression de travail une fois établie, on entretient la décharge électrique entre deux électrodes, à la fréquence de 2 KHz, et le film du polymère se dépose à la surface de ces dernières (fig. 1). L'électrode inférieure est constituée d'un substrat qui est, suivant les études, soit une lame de verre métallisée soit une plaquette de silicium.

On a étudié les paramètres qui régissent la cinétique de croissance de ces films, c'est-à-dire :

- la pression du monomère,
- la température du substrat,
- la durée de la décharge électrique,

- sa densité de courant et sa fréquence. Mais les mécanismes chimiques qui sont à l'origine de la polymérisation dans la décharge électrique ont aussi été analysés.

Caractérisation du matériau et applications

Les propriétés électriques du polymère ainsi déposé étant différentes de celles du polystyrène élaboré de façon classique, il a fallu caractériser par spectrographie infrarouge et analyse chimique, le matériau obtenu. Les expériences ont montré que ce type de polystyrène était très réticulé et oxydé. Il est apparu en outre que la pression de monomère influait sur ses propriétés chimiques.

Le polymère obtenu présente un indice de réfraction de 1,708, une rigidité diélectrique de $5 \cdot 10^6$ V/cm, une permittivité relative de 2,3 et un facteur de pertes diélectriques de 0,005 (ces deux dernières valeurs correspondant à une fréquence de mesure de 1000 Hz).

Ces films peuvent être déposés suivant des épaisseurs qui vont de 150 Å à 5 µm. Ce sont les couches les plus min-

ces que l'on sache obtenir dans ce domaine et, même pour les plus fines d'entre elles, un examen au microscope électronique à balayage montre qu'elles sont exemptes de défauts.

De tels dépôts ont été réalisés sur les dispositifs à semi-conducteurs (diodes et transistors) afin de les protéger contre l'humidité et, plus généralement, pour les "passiver" c'est-à-dire empêcher que leurs caractéristiques électriques ne varient au cours du temps. A l'heure actuelle, les dispositifs à semi-conducteur sont pour la plupart passivés par une couche de silice qui présente quelques défauts : perméabilité à certaines impuretés et surtout sensibilité aux radiations, ce qui devient rédhibitoire pour des dispositifs destinés à l'espace. Il existe également des composants de puissance, diodes et thyristors, fabriqués selon la technique "mesa" dont la passivation est obtenue au moyen de silicones déposés par projection.

Le dépôt de couches minces de polymère par décharge électrique permet quant à lui d'opérer la passivation d'une manière plus rationnelle et plus économique en ce qui concerne les structures

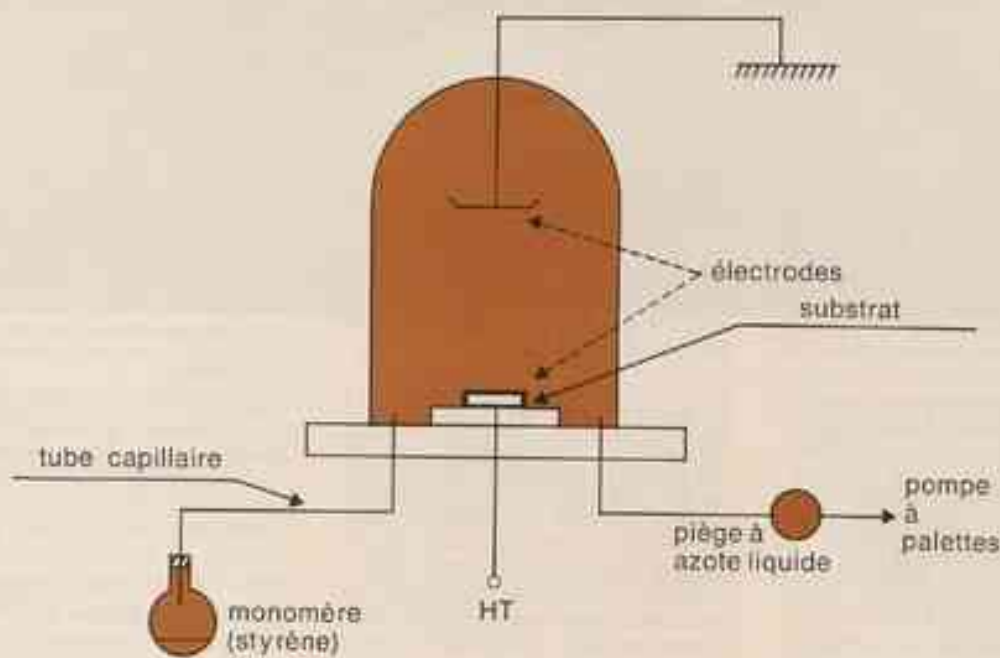
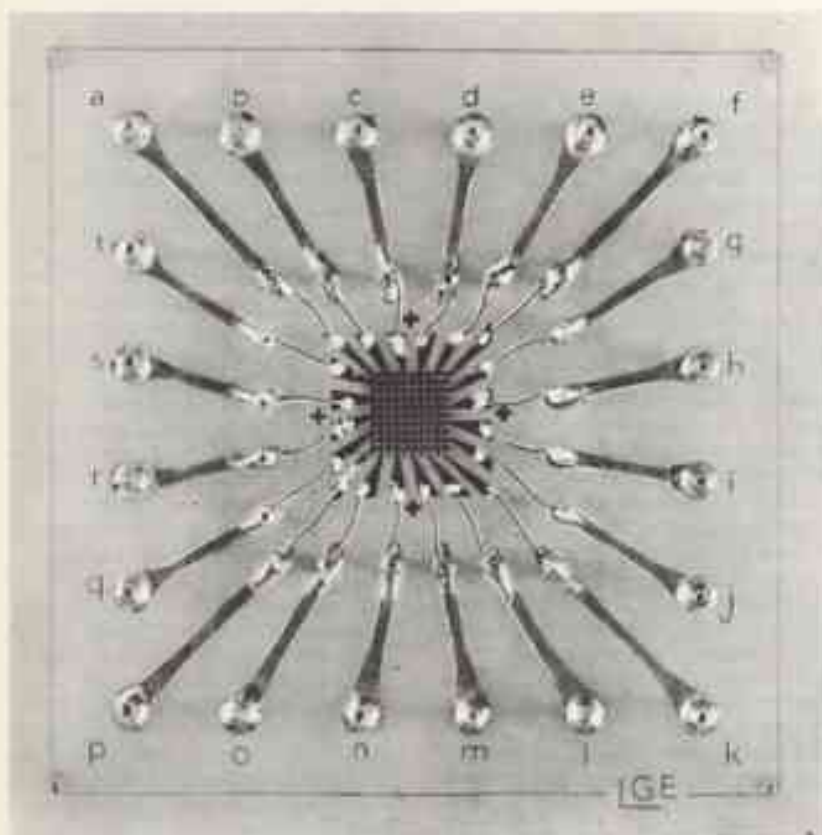


Schéma du dispositif de dépôt des couches minces de polymère.

- Laboratoire de Génie Électrique - Toulouse



Cliché du prototype d'une matrice de cent points mémoires. Chaque structure métal-polymère-métal, représentant un point mémoire, a une surface de $0,5 \times 0,3 \text{ mm}^2$.

"mésa". D'autre part il s'est avéré que les dispositifs ainsi traités, et exposés à des rayons X (tension d'accélération 150 KeV, 10 ma de courant d'anode) ou à des bombardements électroniques (tension d'accélération 50 KeV, flux 10^7 à $10^{11} \text{ e cm}^{-2} \text{ s}^{-1}$) ont une meilleure tenue à ces irradiations que les éléments passivés traditionnellement avec de la silice. On peut noter à ce propos que les doses de radiations utilisées dans ces expériences sont équivalentes à la dose totale reçue par les composants au cours d'une mission spatiale. La solution apportée au problème de la passivation, outre son originalité et son efficacité, présente surtout l'avantage d'une mise en œuvre beaucoup plus facile que celle des palliatifs utilisés jusqu'à ce jour, tel que le "durcissement" de la silice ou l'emploi d'alumine.

La seconde application de ces couches minces concerne la réalisation de micro-condensateurs à capacité volumique élevée. Actuellement les feuilles de diélectrique les plus minces présentent

des épaisseurs de l'ordre de $2 \mu\text{m}$ si bien qu'en utilisant un film de polymère de $0,2 \mu\text{m}$, on multiplie par un facteur 10 la capacité par unité de surface du condensateur. De plus, pour la même épaisseur totale de $2 \mu\text{m}$, on peut empiler dix couches élémentaires qui, mises en parallèle, conduisent à une capacité totale multipliée par un facteur 100, d'où un accroissement appréciable de la capacité volumique.

Une troisième application a donné lieu au dépôt d'un brevet. Les structures métal-polymère-métal obtenues avec ces films ont la propriété de présenter un état isolant et un état conducteur stables. Lorsque la structure vient d'être élaborée, elle présente une résistance d'environ $10^8 \Omega$ pour 2000 Å d'épaisseur. Si l'on applique une tension continue croissante entre ses électrodes, on atteint une tension de seuil V_s pour laquelle la structure bascule rapidement (10^{-8} s environ) dans un état conducteur caractérisé par une résistance d'environ 5Ω . Il suffit en-

suite d'une impulsion de courant pour ramener la structure à l'état isolant. Ces deux états sont stables et surtout ne nécessitent pas une tension de maintien pour rester permanents. Ces basculements isolant-conducteur peuvent être opérés environ 2000 fois avant que le dispositif ne soit altéré. Cependant entre le temps d'application de la tension V_s et la commutation, il s'écoule une durée d'environ 10^{-7} s . Ce retard à la commutation et le nombre relativement restreint de basculements permis font que ces dispositifs trouvent surtout leur intérêt dans la réalisation de mémoires mortes non volatiles et reprogrammables un certain nombre de fois. La photo donne une vue d'un prototype d'une matrice de 100 points mémoires.

Il est intéressant de mentionner également une autre application qui est à l'étude depuis peu en France et dans un grand laboratoire des Etats-Unis : la réalisation, avec ce type de dépôts, de guides de lumière pour l'optique intégrée. En effet, ces films sont transparents et leur indice de réfraction étant supérieur à celui du verre servant de substrat, le rayon d'un laser qui le pénètre à partir d'un prisme (fig. 2) se propage le long de ces couches et peut être ensuite restitué vers l'extérieur à l'aide d'un autre prisme. Les films sont donc intéressants pour les circuits intégrés optoélectroniques qui seront employés dans les futurs systèmes de communication par laser.



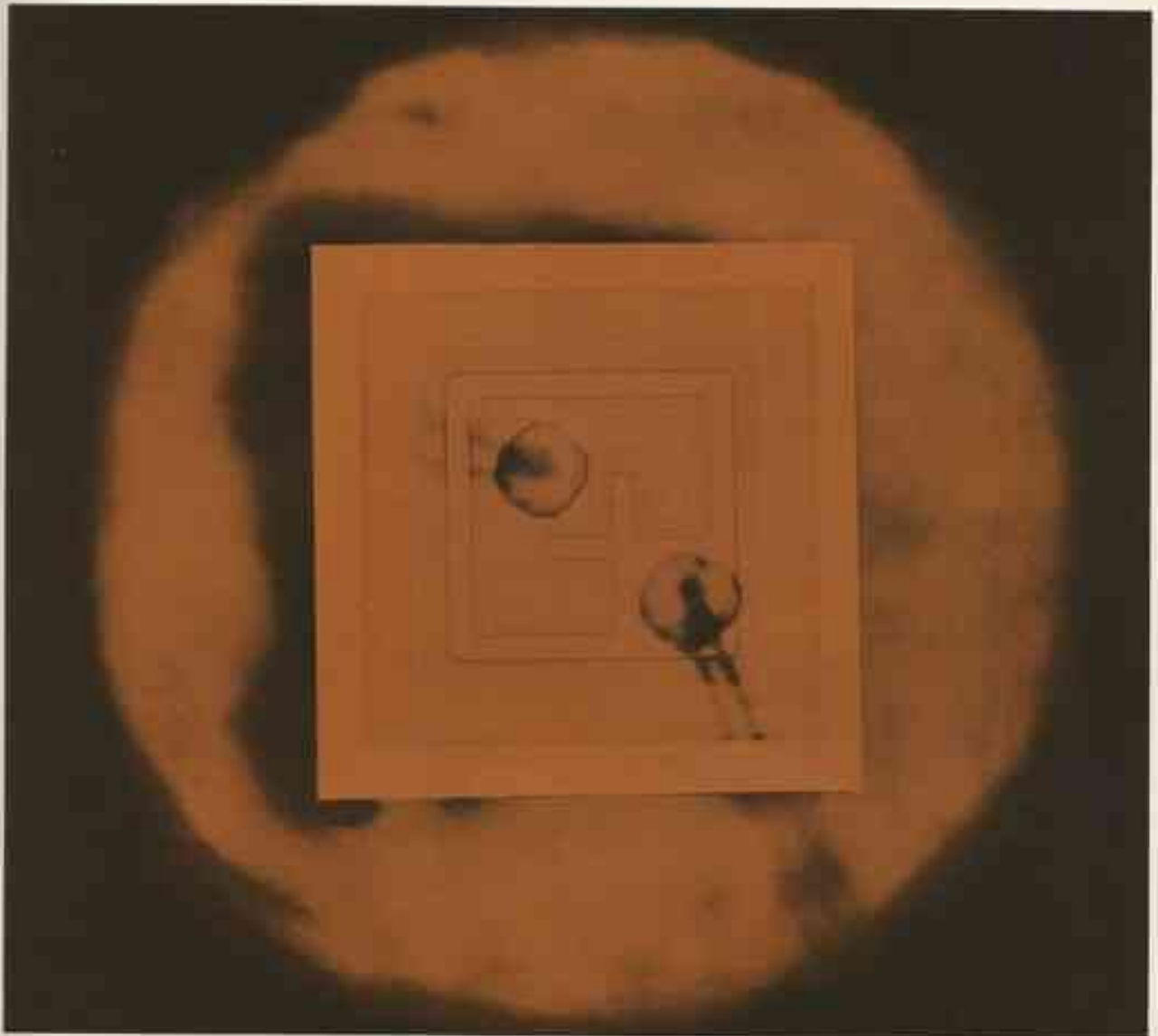
Schéma montrant le principe du montage utilisé lorsque la couche de polymère sert de guide de lumière.

Le faisceau du laser pénètre dans le film diélectrique à l'aide d'un couplage par prisme et en sort de la même façon. Entre les deux prismes, la lumière est véhiculée par la couche mince de polymère.

Enfin il existe un certain nombre d'applications qui ne sont pas développées en France mais envisagées à l'étranger :

- utilisation de ces couches pour les revêtements anticorrosifs,
- réalisation de membranes pour la désalinisation de l'eau de mer.

modèles de composants électroniques semi-conducteurs



Un circuit électronique (Photo L.A.S.).

Depuis leur apparition en 1948, les composants semi-conducteurs actifs ont connu un développement considérable. Hormis toute considération économique, un tel essor est le résultat de la mise au point de nouveaux procédés technologiques, notamment la technologie planar, ainsi que des progrès continus réalisés dans le contrôle de ces procédés.

De ce fait, outre l'amélioration des performances des composants discrets, on a pu concevoir et réaliser des cir-

cuits intégrés dont la densité atteint aujourd'hui plusieurs milliers de composants interconnectés sur un même substrat de semi-conducteur. On se trouve donc en présence de macro-

- Laboratoire d'Automatique et d'Analyse des Systèmes - Toulouse
- Institut d'Electronique Fondamentale - Paris XI
- Centre d'Etude d'Electronique des Solides - Université des Sciences et Techniques du Langedoc
- Centre de Recherche sur les Propriétés Hyperfréquences des milieux condensés - Université des Sciences et Techniques - Lille

composants qui réalisent des fonctions complexes.

Cette situation nouvelle a entraîné une mutation profonde dans la conception et l'analyse des propriétés des composants élémentaires (transistors bipolaires, transistors à effet de champ, diodes, etc...). En effet dans l'optique des fonctions intégrées, la conception des circuits de l'électronique, assistée ou non par ordinateur, a imposé la connaissance préalable des propriétés des éléments constitutifs.

Mais à cette connaissance vient s'ajouter une exigence de "précision" dans la modélisation du comportement électrique des composants. Outre les propriétés intrinsèques des éléments, il devient nécessaire de tenir compte des phénomènes parasites. L'importance relative de ceux-ci est d'autant plus grande que les progrès technologiques ont conduit à une amélioration des propriétés intrinsèques par un meilleur contrôle des différentes étapes de fabrication.

Deux voies d'approche

Pour satisfaire à ces critères, on peut envisager deux voies d'approche pour élaborer les modèles des composants.

On peut effectuer une identification mathématique des propriétés observées expérimentalement sans analyser leur origine physique (technique de "la boîte noire"). Cette méthode dite déductive, présente un inconvénient majeur : toute "entrée nouvelle" par rapport aux conditions d'identification (modification technologique, influence de contraintes, ...) a des conséquences que le modèle est incapable de prévoir, ce qui impose à chaque fois une nouvelle séquence d'identification.

Dans l'autre méthode, dite inductive, on procède à l'identification et à l'analyse de chaque mécanisme physique élémentaire, intrinsèque ou parasite, impliqué dans le comportement du dispositif. Le modèle est alors l'image synthétique de l'ensemble des résultats obtenus. Elaboré sur des bases physiques, il permet de mettre en évidence, d'analyser et de prévoir l'action des contraintes (température et irradiations par exemple) ; il constitue un outil particulièrement utile pour étudier l'influence des conditions technologiques d'élaboration sur les propriétés électriques résultantes.

Deux modèles de composants

C'est dans cette dernière optique qu'ont été abordées les études relatives à la modélisation entreprises au Labora-

toire d'Automatique et d'Analyse des Systèmes.

En ce qui concerne le transistor bipolaire, les recherches ont abouti à la proposition de modèles analytiques et compacts : IBIS, en régime statique ou continu et BIRD en régime dynamique ou alternatif. Rappelons que les transistors bipolaires sont des composants qui comportent trois couches de semi-conducteur dont les porteurs majoritaires sont alternativement des électrons et des trous et dans lesquels l'effet d'amplification est obtenu par injection de porteurs minoritaires dans la base, leur collection étant réalisée par le collecteur. Les modèles IBIS et BIRD tiennent compte non seulement de l'effet transistor proprement dit mais de la nature bidimensionnelle des phénomènes de conduction dans la zone active de base. Ils tiennent compte également de l'existence de zones latérales et superficielles, du comportement résistif de la région du collecteur et des éléments parasites dus à l'embase et au boîtier.

Ces modèles sont aptes à traduire toutes les particularités de fonctionnement des transistors bipolaires constatées expérimentalement ; il convient de citer notamment les variations, avec les conditions de polarisation et la fréquence, du gain en courant et de la transadmittance.

Voyons maintenant les modèles qui ont été choisis pour les transistors à effet de champ à grille isolée. Dans ces composants l'effet d'amplification est obtenu par la modulation au moyen d'un champ électrique de la densité des porteurs qui sont localisés dans un canal situé près d'une interface isolant-semi-conducteur. Les travaux effectués ont permis d'élaborer des modèles qui rendent compte des propriétés statiques et dynamiques en prenant en considération les variations de mobilité des porteurs sous l'influence des champs électriques longitudinaux et transversaux. De tels modèles tiennent compte aussi de l'ensemble des éléments parasites notamment les capacités dues aux différentes configurations du composant.

Les schémas obtenus simulent fidèlement le comportement du composant,

en fonction des conditions de polarisation et de la fréquence, et en particulier les variations en tension grille de la transconductance et celles en fréquence de l'admittance de drain.

Tous ces modèles sont caractérisés par des paramètres indépendants des conditions de polarisation qui sont accessibles par des mesures simples et sont directement liés à des phases précises du processus technologique d'élaboration. Chaque paramètre est fonction en particulier des dimensions géométriques des masques. Par ailleurs, ils permettent de prévoir le comportement thermique des différents types de composants.

Applications

Ces modèles ont fait l'objet de deux types d'applications.

Ils ont été utilisés pour étudier les origines des dispersions de fabrication permettant ainsi de caractériser une technologie à partir de mesures électriques. On a en particulier mis en évidence l'importance prépondérante des dispersions liées aux procédés de passivation par rapport à celles provoquées par les processus de diffusion.

Sur le plan de la conception des composants, on a pu définir avec précision l'influence des dimensions géométriques des masques et déterminer ainsi les facteurs d'échelle à appliquer entre les différents composants fabriqués à partir d'un même processus technologique.

D'autre part, dans le cadre d'études plus fondamentales, ils ont servi d'outil d'analyse pour mettre en évidence l'action des radiations ionisantes (création de charges fixes dans les couches de silice, modifications des propriétés d'interface, ...). Ils ont permis de mener à bien les études sur la localisation et l'origine des sources de bruit de fond.

Enfin, dans le cas des transistors bipolaires, ils servent aussi de base à l'étude du comportement des structures de puissance pour lesquelles il est nécessaire de tenir compte de la répartition non homogène de la température dans le cristal de semi-conducteur.

l'automatique... une des disciplines types des sciences pour l'ingénieur

identification par filtrage non linéaire

Si on imagine un processus quelconque (avion en vol, colonne de distillation, réacteur chimique...), un modèle est un instrument susceptible de remplacer le processus et qui permet de pratiquer certaines expériences ou d'effectuer diverses études. Un système physique (simulateur) sera considéré comme un modèle, s'il présente, dans un certain domaine, un comportement analogue à celui de l'objet étudié. A la limite, le modèle peut être défini comme un objet mathématique (système d'équations par exemple), indépendamment du système physique qui permettrait de le réaliser. Au sens où l'on entend généralement l'automatique, l'identification est la recherche d'un modèle mathématique, susceptible de représenter en particulier le comportement dynamique d'un système à contrôler. Un modèle est adéquat s'il permet de prévoir l'évolution dans le temps de certaines variables dépendantes du système (sorties), en fonction de l'évolution des variables indépendantes (entrées).

Comment élabore-t-on un tel modèle ? L'identification d'un processus donne lieu à l'effectue à partir de deux types d'informations :

- des informations primaires, issues de l'analyse physico-mathématique de chaque élément de base du système. Le modèle mathématique peut alors être constitué de la réunion de toutes les équations élémentaires accompagnées des valeurs numériques correspondantes ;
- des informations secondaires (ou de

travail) constituées par des mesures effectuées sur les entrées-sorties du processus en régime dynamique.

Certaines informations primaires sont indispensables, au moins pour définir la forme générale du modèle (tel type d'équations différentielles ou récurrentes), mais c'est généralement à partir des informations secondaires que l'on détermine les valeurs numériques des coefficients (ou paramètres) de ces équations, de façon à obtenir la similitude entre l'objet et le modèle.

On envisagera ici ce deuxième aspect de l'identification, c'est-à-dire seulement l'estimation des paramètres à partir des mesures d'entrée-sortie.

Estimation des paramètres

Il est impossible de parler d'identification sans dire d'abord un mot de la méthode dite du modèle. Il s'agit en fait d'une méthode générale qui englobe de nombreuses techniques particulières. Elle consiste à réaliser, par exemple à l'aide d'un calculateur numérique ou analogique, un système simulateur dont les paramètres sont ajustables. On adapte alors les paramètres de ce modèle jusqu'à ce que la sortie du modèle soit identique (à un certain degré d'approximation près) à la sortie de l'objet et ceci pour le même signal d'entrée (figure 1). La détermination d'une procédure générale d'ajustement des para-

mètres n'est pas un problème simple à cause surtout de l'aspect dynamique des phénomènes. Pour le résoudre, il faut faire appel à des théories comme celles de la sensibilité, de la programmation non linéaire, de l'hyperstabilité...

La méthode du modèle est particulièrement bien adaptée à ce qu'on peut appeler le cas réel, où l'hypothèse de l'identité de la structure mathématique de l'objet et du modèle peut être hasardeuse, ou parfois même formellement contredite lorsqu'on recherche un modèle simple pour un système complexe. Néanmoins, dans certains cas, on peut pousser assez loin les hypothèses sur le processus à identifier, celles-ci pouvant concerner jusqu'aux propriétés statistiques des bruits de mesure, des perturbations, et même de l'évolution des paramètres.

Un exemple caractéristique est celui d'un réacteur nucléaire pour lequel les lois de la physique statistique permettent d'écrire les équations générales de fonctionnement, équations différentielles stochastiques par exemple. Les paramètres de ces équations sont alors déterminés, grâce à l'identification à partir des mesures d'entrée-sortie sur le réacteur en fonctionnement. Dans ce cas, l'identification relève de la théorie de l'estimation statistique et c'est ici que peut s'introduire l'identification par filtrage non linéaire.

Filtrage non linéaire

Le terme de "filtrage" est équivalent à l'estimation de l'état d'un système décrit par des équations différentielles. Il doit donc permettre de suivre les variations de la "trajectoire" de l'état à l'inverse de l'estimation statistique qui porte sur des constantes. L'information qui permet cette estimation est fournie uniquement par les entrées et les sorties du système, ces dernières étant considérées comme une image indirecte et perturbée de son état (figure 2). Non seulement le système présente les caractères aléatoires désignés sous le nom de perturbations ou erreurs mais encore certains paramètres indispensables à la représentation même du système sont inconnus.

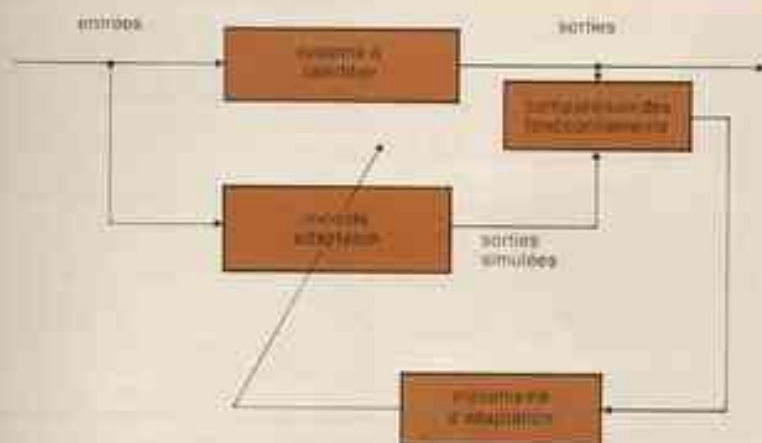


schéma de principe de l'identification d'un système par la méthode du modèle

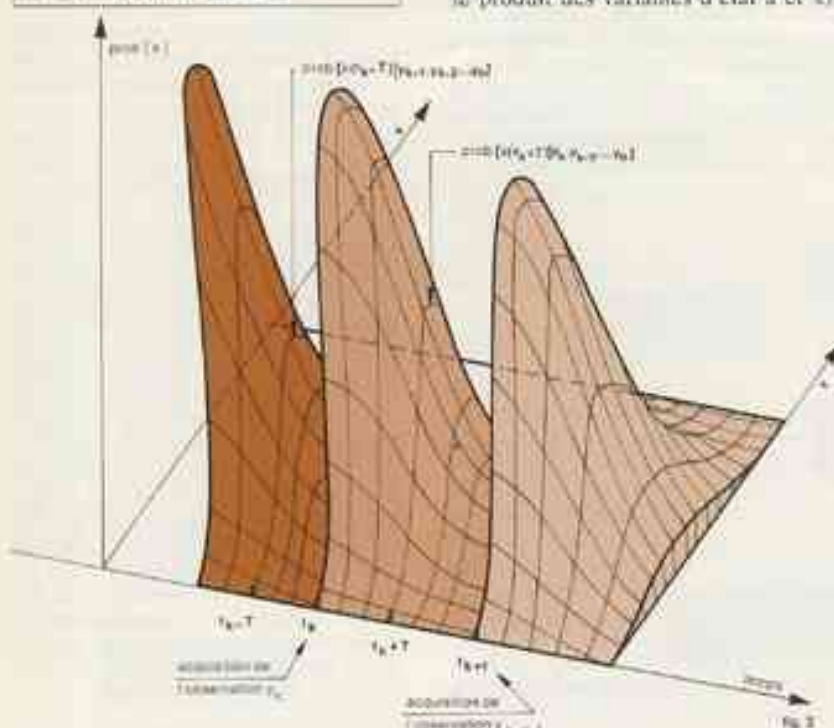
fig. 1

A l'aide d'un calculateur numérique ou analogique, on réalise un modèle dont les paramètres sont adaptés jusqu'à ce que son fonctionnement soit équivalent à celui du système. De nombreuses techniques sont utilisées pour réaliser cette adaptation parmi lesquelles se place celle du filtrage non linéaire.

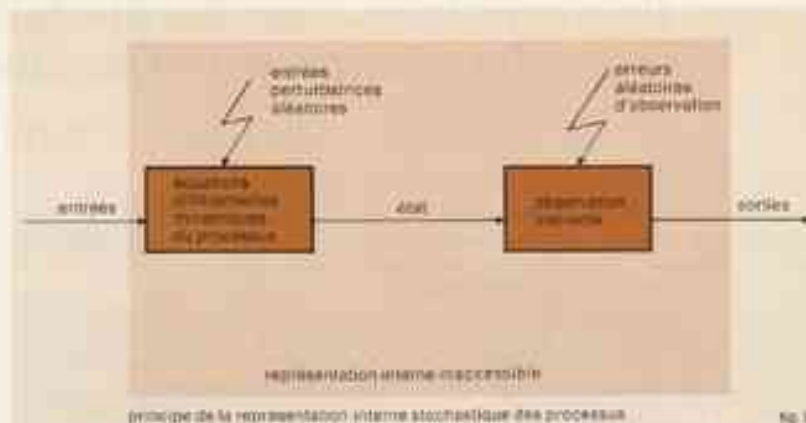
• Laboratoire d'Automatique et d'Analyse des Systèmes - Toulouse.
• Laboratoire d'Automatique - de l'École Nationale Supérieure de Mécanique - Nantes.

Pour appliquer la théorie du filtrage à l'identification, une solution maintenant classique consiste à considérer les paramètres à identifier comme des états supplémentaires dont on vient augmenter le vecteur d'état proprement dit. Pour cela, il faut écrire la loi d'évolution des paramètres sous forme d'équation d'état, ce qui est trivial au moins dans le cas stationnaire où il suffit d'écrire que les variations des paramètres sont nulles.

$\frac{dx}{dt} = ax + bu + v$ (1) - u = entrée ou commande - v = entrée aléatoire ou perturbation - x = variable d'état - a et b = paramètres	
$y = hx + e$ (2) - h = paramètre - e = erreur aléatoire - y = sortie	
$\frac{da}{dt} = 0$ (3) - Equation d'état à ajouter à (1) - $\begin{bmatrix} x \\ a \end{bmatrix}$ = nouveau vecteur d'état augmenté	



À tout instant la connaissance imparfaite de l'état du système est donnée par une densité de probabilité. En absence d'observation, cette connaissance se dilue et la densité s'aplatit ou se disperse. L'acquisition d'une nouvelle observation remet à jour la connaissance et regroupe la probabilité.



Le modèle adopté par la théorie des systèmes pour représenter leur évolution dynamique fait apparaître les notions d'entrées, d'état et de sorties : ce schéma est à rattacher à celui d'une interrogation de la nature à l'aide d'actions expérimentales qui provoquent des effets, observés directement, qui reflètent imparfaitement des caractères inaccessibles qui déterminent son état.

Considérons par exemple le système décrit par l'équation différentielle linéaire (1) et par l'équation d'observation (2). Si a est un paramètre inconnu à estimer et supposé constant, on ajoute à (1) l'équation (3), ce qui revient à considérer un nouveau système dynamique formé de (1) et (3), muni de la même équation d'observation (2). Le résultat de cette "immersion" des paramètres dans l'état aboutit toujours, même dans le cas exposé qui est extrêmement simple, à un système d'équations dynamiques finales de dimension plus élevée et qui présente obligatoirement des non linéarités (ici le produit des variables d'état a et x).

L'identification des paramètres se présente alors comme un filtrage appliqué au vecteur d'état augmenté.

Cette méthode semble pouvoir s'appliquer à une large classe de systèmes, non linéaires ou non stationnaires. De plus, elle permet d'effectuer une identification continue du processus grâce à de faibles perturbations qui affectent peu son fonctionnement normal.

Malheureusement, les résultats actuellement disponibles de façon générale et compacte concernent seulement l'estimation de l'état de systèmes linéaires (équation (1) seule par exemple) à paramètres constants (filtre de Wiener) ou à paramètres variables de façon connue (filtre de Kalman-Bucy).

La théorie du filtrage non linéaire, c'est-à-dire en général de l'estimation de l'état de systèmes décrits par des systèmes d'équations différentielles non linéaires, est assez ardue et la solution théorique exacte est encore très récente. Elle se présente sous forme d'itérations des densités de probabilités conditionnelles remises à jour après chaque nouvelle observation, en tenant compte de l'évolution (ou diffusion) de cette probabilité introduite par la dynamique du processus. La figure 3 représente trois instants choisis avant, pendant et après l'acquisition d'une nouvelle observation (sortie). On peut ainsi observer :

- La diffusion ou dispersion de la probabilité entre deux instants quelconques.
- La concentration de la probabilité lors de la prise en considération d'une nouvelle observation.

Il en résulte que les équations de filtrage sont, dans le cas général, des équations de diffusion aux dérivées partielles du type Fokker-Planck, pour lesquelles il est en général impossible de trouver une solution numérique exacte dès que la dimension de l'état dépasse quelques unités.

Les travaux du L.A.A.S. et de l'E.N.S.M. sur l'identification par filtrage non linéaire

Au Laboratoire d'Automatique de l'E.N.S.M., l'attention s'est portée sur l'identification de modèles discrets par filtrage non linéaire approché. Certaines difficultés théoriques, liées aux équations différentielles stochastiques, disparaissent alors, ce qui a permis d'axer les recherches sur les conditions d'application à des processus de dimension relativement élevée.

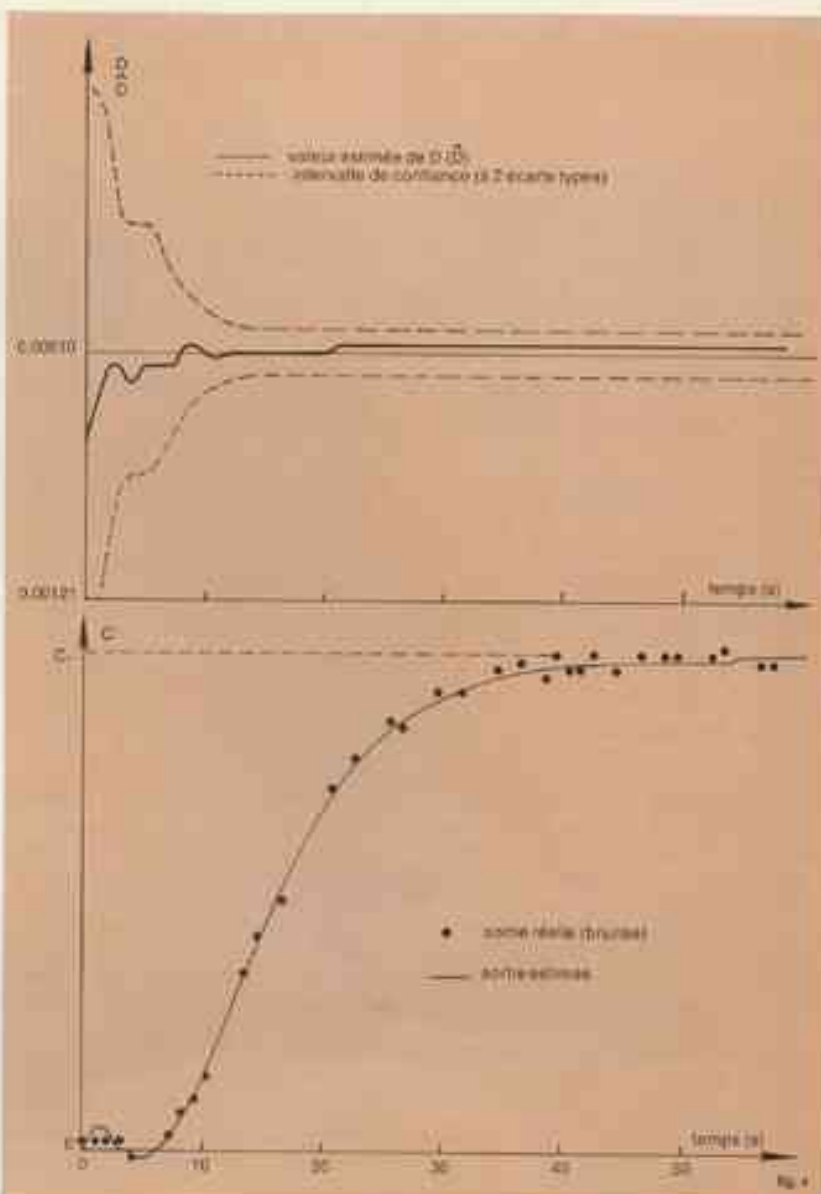
La méthode de filtrage continu-discret et celle du filtre approché entièrement discret ont permis de mettre au point des ensembles de sous-programmes de base avec lesquels on peut réaliser, avec un minimum d'efforts, les filtres correspondant à des problèmes relativement divers.

L'application a été faite à des processus tels que les réacteurs chimiques (Rhône Progil), ou centrale thermique (E.D.F.). On peut noter que les méthodes développées à l'E.N.S.M. ont été transmises à la S.N.E.C.M.A. qui les utilise systématiquement pour l'identification de turbo-réacteurs (systèmes non linéaires où le nombre de paramètres à identifier simultanément peut atteindre dix à vingt).

Des chercheurs du L.A.A.S. se sont préoccupés du problème de la résolution exacte des équations du filtrage, au moins dans le cas de systèmes d'ordre faible, en utilisant notamment les solutions explicites données par le théorème de représentation de Bucy.

D'autres travaux concernent l'approximation de la densité de probabilité conditionnelle à l'aide de développements fonctionnels limités. Cette méthode est particulièrement adaptée au cas où l'on considère une évolution dynamique continue, alors que les observations sont échantillonnées.

Les applications à ce jour concernent le traitement de signaux d'entrée-sortie d'un réacteur nucléaire (C.E.A.), d'un réacteur chimique (S.N.P.A.) et d'un réacteur chimique pilote (L.A.A.S.), en vue de leur identification.

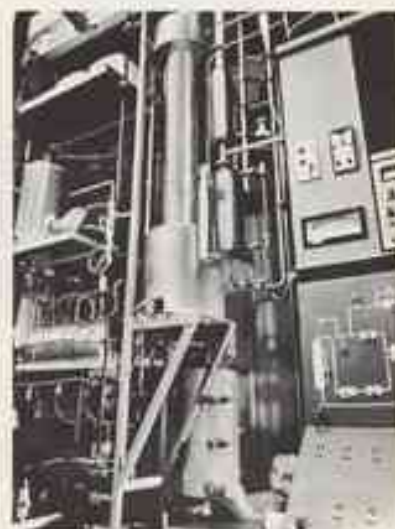


Courbes montrant l'évolution du paramètre estimé avec son intervalle de confiance ainsi que la comparaison des sorties du modèle ainsi identifié et du système réel.

Cette dernière application peut fournir un exemple illustratif des résultats auxquels conduit l'estimation d'un paramètre par la méthode du filtrage non linéaire. A partir des mesures faites sur l'unité pilote (photo), la figure 4 reproduit les courbes d'évolution des paramètres estimés avec son intervalle de confiance ainsi que la comparaison des sorties du modèle ainsi identifié et du système réel.

Ce survol rapide de l'activité de quelques laboratoires à propos du filtrage non linéaire ne prétend pas donner une idée générale des recherches en France sur l'identification, mais peut permettre au lecteur de se faire une opinion sur l'intérêt de ces problèmes et du caractère des recherches qui en découlent.

Unité Pilote chimique (Photo Yan).



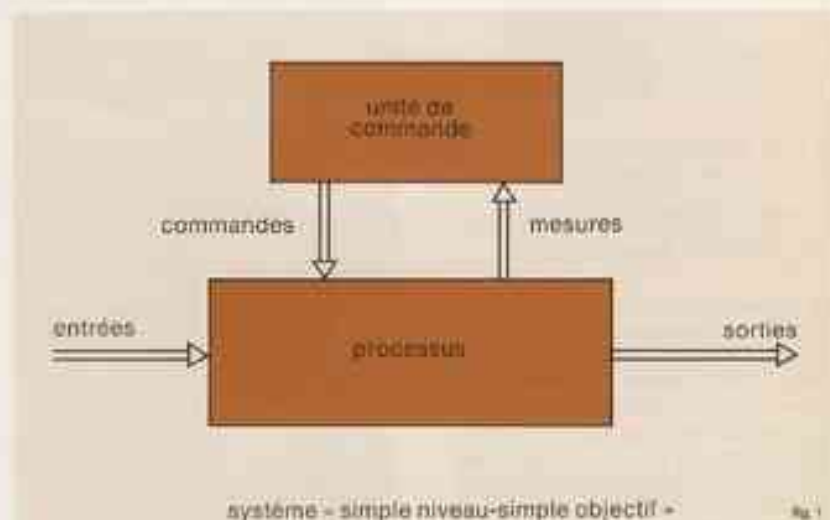
commande des systèmes complexes

Le problème de la commande optimale d'un processus consiste à déterminer et à appliquer la politique (les commandes en fonction du temps), réalisant la consigne (la tâche fixée) et extrémalisant un critère de fonctionnement, tout en satisfaisant les contraintes. Cette notion d'optimisation est à la base de nombreux travaux en Automatique; ceci a conduit à la proposition de méthodes de commande optimale récemment développées mais qui, pour certaines, découlent de théories classiques. Ces méthodes peuvent se classer en trois catégories :

- programmation linéaire et non linéaire
- programmation dynamique
- méthodes variationnelles (calcul classique des variations et principe du maximum).

Limitation de ces méthodes

Si les techniques de l'Automatique mentionnées ci-dessus ont été longtemps appliquées à la commande de systèmes relativement simples, elles s'étendent maintenant de plus en plus à la commande de processus complexes, de grande dimension (processus industriels ou économique-industriels). Cette mutation étroitement liée à la généralisation de l'introduction des calculateurs dans les chaînes de commande, entraîne à l'heure actuelle une évolution profonde des problèmes de commande tant au niveau de la complexité des modèles entrant en jeu qu'au niveau de la complexité des fonctions de commande à assurer.



Dans une approche globale d'un problème de commande, le système de commande alors unique réalise les mesures sur le processus considéré comme un tout, et élabore, à partir de cette information, une commande qu'il applique au processus.

L'application directe, dans une approche globale, des méthodes classiques de l'Automatique à la résolution de tels problèmes présente de nombreuses difficultés, notamment de calcul; aujourd'hui de telles méthodes ne sont pas rentables.

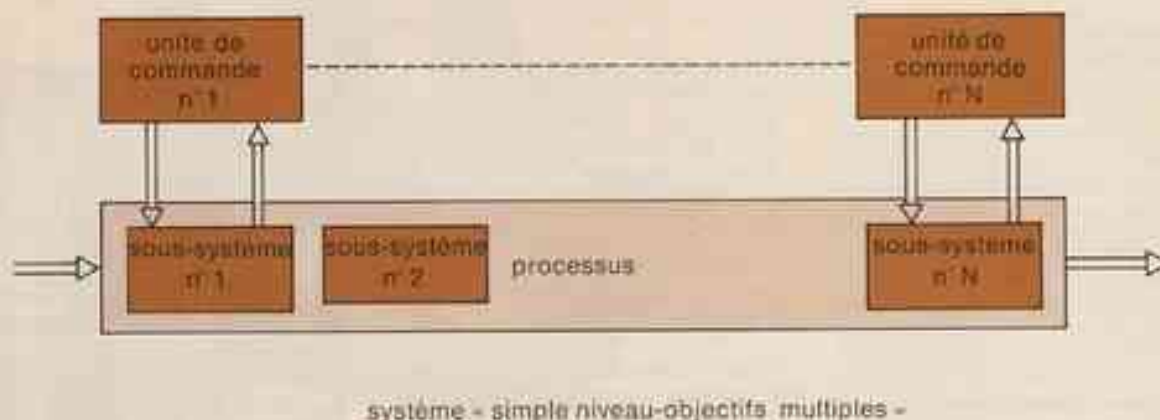
Abandonnant toute approche globale, on a donc cherché à résoudre le problème en le divisant en sous-problèmes plus simples, c'est-à-dire à commander le système à des niveaux de complexité croissante. C'est cette démarche axée

Laboratoire d'Automatique et d'Analyse des Systèmes - Toulouse, Laboratoire d'automatique de l'E.N.S.M. - Nancy.

naturelle que suit la commande hiérarchisée qui, dans l'état actuel des connaissances, s'avère la plus prometteuse pour la commande des systèmes complexes.

Ce nouveau type de commande proposé initialement par Mesarovic (Etats-Unis) vers 1964 a suscité depuis de nombreux travaux notamment aux Etats-Unis, en U.R.S.S. et en POLOGNE. Les premières recherches françaises dans ce domaine datent de 1969 et le nombre des Laboratoires d'Automatique abordant ce sujet va sans cesse croissant.

Quel bilan peut-on dresser aujourd'hui de ces travaux ?



Dans l'hypothèse où le processus est considéré comme composé de sous-systèmes indépendants, chacun de ces sous-processus peut être commandé par une unité de commande locale, ce qui définit une structure "simple niveau-objectifs multiples". Cependant cette hypothèse d'indépendance n'étant jamais rigoureusement satisfaite, la commande obtenue est, en général, sous-optimale.

La commande hiérarchisée

Lors d'une approche globale directe d'un problème de commande, les relations entre le système de commande alors unique et le processus, considéré comme un tout, peuvent être représentées par un schéma dit à "simple niveau - simple objectif" (fig. 1). Nous avons vu précédemment les difficultés qui sont liées à un tel principe de commande lorsque le processus et (ou) la fonction de commande deviennent complexes.

Mais si on peut considérer le processus global comme étant formé de sous-processus indépendants, on définit alors un système de commande dit à "simple niveau - objectifs multiples" (fig. 2). Cette hypothèse d'indépendance n'est pourtant pas physiquement réaliste et la solution obtenue pour le système global est une solution sous-optimale d'autant plus éloignée de la solution réelle donnée par une méthode directe que le couplage entre sous-systèmes est faible. Lorsque les interactions sont fortes, un objectif étant associé à chaque unité de commande, des conflits peuvent apparaître entre ces unités sans qu'aucune d'elles n'ait a priori la possibilité de les résoudre : cette approche n'a alors plus de sens et on doit avoir recours à un deuxième niveau de commande qui tiendra compte de ces interactions et résoudra les conflits en fixant des objectifs appropriés à chaque unité de commande.

On est donc conduit logiquement à définir des systèmes de commande à "plusieurs niveaux - plusieurs objectifs" (fig. 3). Les unités de commande sont réparties sur deux ou plusieurs niveaux suivant une hiérarchie et certaines n'ont qu'un accès indirect au processus à commander. Ces unités traitent les informations reçues pour commander ensuite d'autres unités qui leur sont inférieures dans la hiérarchie et renvoient le résultat de leur action vers les niveaux supérieurs.

Les unités de commande ont des objectifs différents, indépendants entre eux, et inconnus des autres unités d'un même niveau : ces unités peuvent être en conflit vis-à-vis de l'objectif global mais ici il existe des unités supérieures appelées coordonnateurs pour résoudre les conflits et arriver à la solution. Deux notions fondamentales sont donc à la base de l'élaboration des structures à

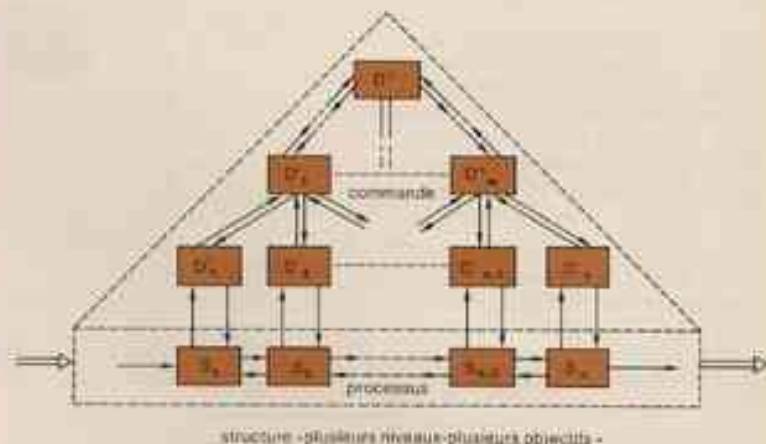


fig. 3

Pour tenir compte des interactions entre les sous-processus et pour résoudre les conflits entre les unités de commande d'un niveau donné, il est nécessaire d'ajouter des niveaux supérieurs de commande, ce qui conduit à une structure hiérarchisée utilisant les notions de division du travail et de coordination des tâches.

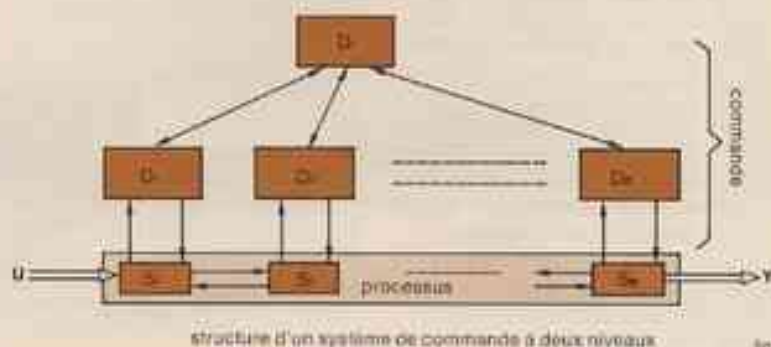


fig. 4

Dans la famille des systèmes de commande hiérarchisée, un système à deux niveaux est particulièrement important : il permet de concevoir, par construction modulaire, des structures plus générales et il met en évidence les problèmes caractéristiques de la commande hiérarchisée (division du travail, transfert d'informations entre niveaux, coordination).

plusieurs niveaux, ce sont :

- la répartition des tâches (division du travail)
- et la coordination, la première nécessitant la seconde.

Face à un processus complexe sur lequel on doit réaliser une fonction de commande complexe, cette division du travail pourra être réalisée suivant deux voies simultanément :

- faisant intervenir un système de couplage, le processus peut être divisé en sous-processus plus simples commandés par des unités de commande locales dont les actions sont coordonnées par des niveaux supérieurs de commande. C'est la division horizontale du travail basée sur la complexité du processus.
- la fonction de commande peut être di-

visée en fonctions plus simples suivant une hiérarchie. C'est la division verticale du travail basée sur la complexité de la fonction de commande.

Dans la classe des systèmes hiérarchisés, un système de commande avec une seule unité de commande au deuxième niveau (fig. 4) est particulièrement important. Il permet notamment de mettre en évidence les problèmes caractéristiques de ce type de commande, en particulier les problèmes fondamentaux de division du travail, transfert d'information entre niveaux, coordination qui peuvent être étudiés sur ce cas particulier. C'est pourquoi la majorité des travaux en commande hiérarchisée portent sur ces structures à deux niveaux.

Trois modes de coordination

La commande hiérarchisée se propose, en définitive, de décomposer un problème global P (dont le but est de satisfaire un objectif global associé au système), en un certain nombre de sous-problèmes P_i dont on recherche la solution localement, de façon à satisfaire : $Sol[P_1 \dots P_n] \Rightarrow Sol[P]$

En fait, une telle relation ne peut être satisfaite à cause des interactions et des conflits éventuels entre les P_i . Il est donc nécessaire d'introduire un "vecteur d'intervention" ou "paramètre de coordination" et de remplacer P_i par $P_i(\lambda)$ pour satisfaire :

$$Sol[P_1(\lambda) \dots P_n(\lambda)] \Rightarrow Sol[P] \\ \lambda = \lambda^*$$

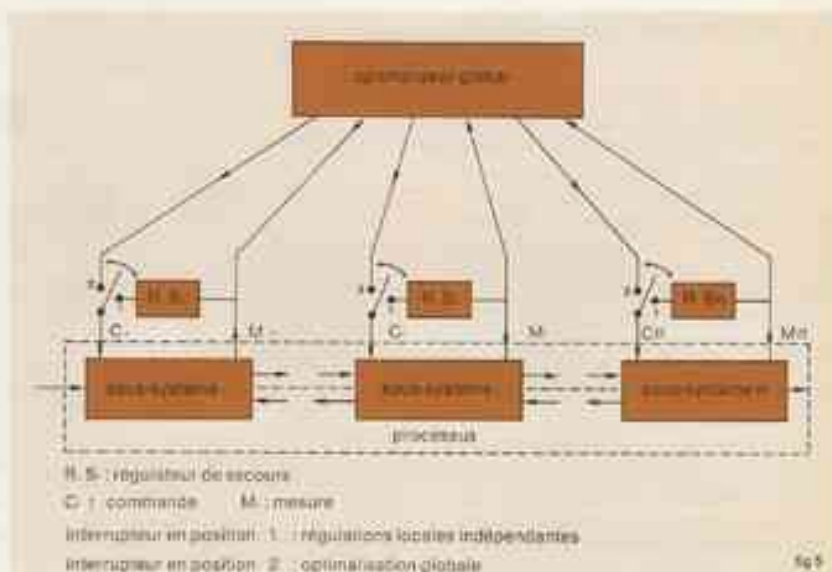
Pour résoudre le problème de la coordination en commande hiérarchisée, il faut faire évoluer $\lambda \rightarrow \lambda^*$, ce dernier conduisant à la solution du problème global. Mais chaque sous-problème est caractérisé par deux fonctions : critère et modèle du sous-système. De ce fait il y a trois types de coordination possibles :

- coordination par action sur le critère ;
- coordination par action sur le modèle ;
- coordination par action sur ces deux fonctions.

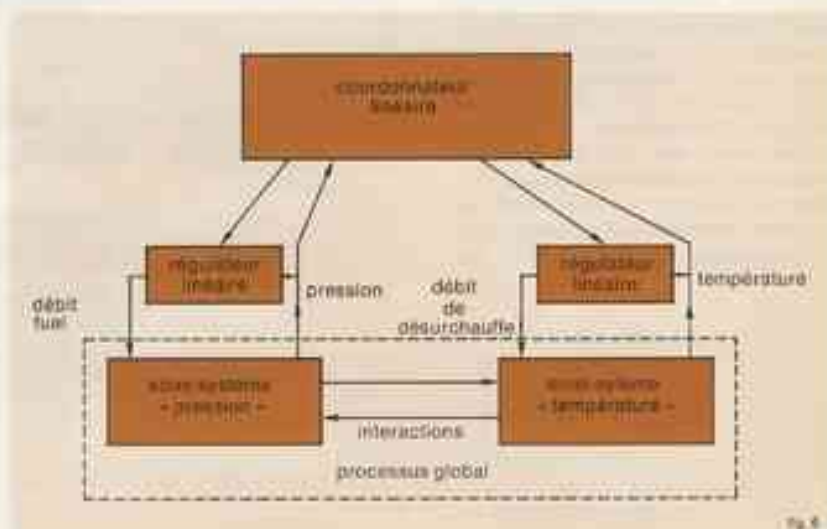
Ce concept de décomposition-coordination utilisé pour la première fois par Dantzig-Wolfe en programmation linéaire a donc été repris ces dernières années dans le contexte de la commande hiérarchisée.

Les chercheurs français se sont attachés essentiellement à l'étude de l'utilisation de la coordination par action sur le critère et le modèle. Leur contribution réside dans la proposition d'algorithmes coordinateurs souvent heuristiques mais très efficaces et dans la résolution de quelques exemples significatifs pour ce type de commande. Cette contribution a permis en particulier de présenter dans une théorie unifiée les trois méthodes possibles de commande optimale à deux niveaux, et ce à partir de la notion de décomposition, en programmation non linéaire pour l'optimisation statique et en calcul des variations pour l'optimisation dynamique.

La notion de décomposition horizontale a donc été largement utilisée et les efforts tendent maintenant vers une utilisation simultanée de deux types de division (horizontale et verticale).



Une pratique industrielle courante consiste à commander un processus composé de sous-systèmes par la juxtaposition de chaînes élémentaires de régulation sur chaque sous-système. Cette pratique subsiste à côté de systèmes de commandes plus évolués et prend le relais de ces derniers lors des démarrages et des arrêts ou lors d'incidents.



Pour contre une structure hiérarchisée permet une approche beaucoup plus réduisante qui intègre les chaînes de régulation locales dans une commande globalement optimale. Ainsi cette approche, appliquée à un générateur de vapeur, fait apparaître des boucles de régulation de température et de pression dont les actions sont coordonnées par un niveau supérieur de commande.

Applications

Actuellement les chercheurs se trouvent dans une phase transitoire très nette. Après les études essentiellement théoriques, ils engagent de sérieuses discussions avec le secteur industriel (industries chimiques, nucléaire, pétrolière, du papier, etc...) en vue d'application aux systèmes complexes de production et

déjà certains contrats de collaboration sont signés.

Ceci explique que pour l'instant les exemples traités en France ne relèvent pas réellement de la commande des systèmes complexes : ils permettent cependant de mettre en évidence soit les possibilités, soit le principe de fonctionnement de ce type de commande.

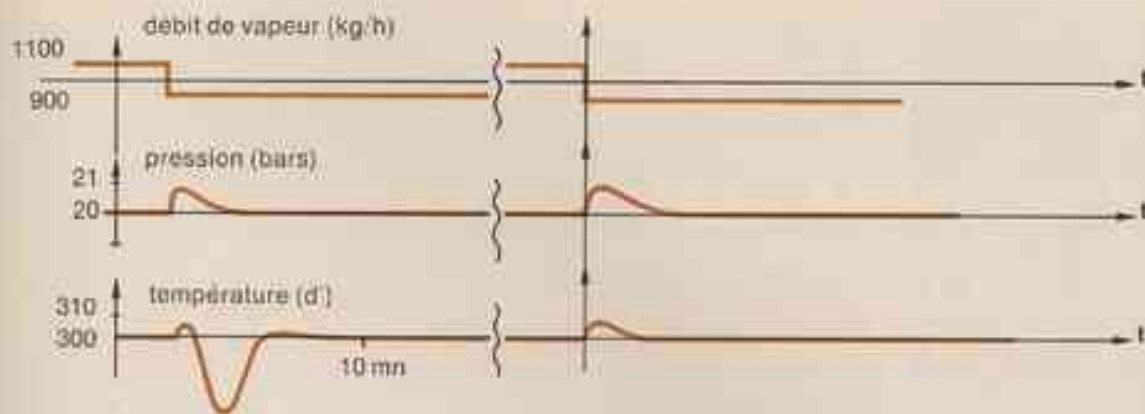


Fig. 7. Fonctionnement en régulation locale :

Lorsqu'une perturbation en échelon par exemple se présente sur le débit de vapeur, elle entraîne une déviation par rapport à la valeur désirée des variables pression et température qui est annulée en un temps convenable par les régulations locales.

Fig. 8. Fonctionnement avec coordination :

La déviation sur les variables pression et température, pour une même perturbation sur le débit de vapeur, est annulée plus rapidement lorsque la coordination entre en jeu. Cette coordination améliore donc considérablement le fonctionnement du système de commande.

qui peut être utilisé dans les domaines suivants :

- optimisation statique (allocation de ressources par exemple) ;
- optimisation dynamique "hors ligne" (calcul des trajectoires optimales) ;
- commande "en ligne" des processus.

Pour citer un exemple, dans un système hydroélectrique, on cherche à obtenir une répartition optimale des énergies de production afin de minimiser les pertes de production et de transmission, pour une demande totale d'énergie fixée. Ce problème, où les sous-systèmes (les centrales) sont fortement couplés (hydrauliquement et électriquement), prend des dimensions importantes lorsqu'il est posé à l'échelon régional ou national ; l'utilisation des techniques de commande hiérarchisée simplifie beaucoup sa résolution.

La commande optimale des systèmes décrits par des équations aux dérivées partielles pose de nombreux problèmes théoriques et de calcul. Une discrétisation dans l'espace ou dans le temps et l'espace a permis d'en résoudre de nombreux par l'utilisation de la commande hiérarchisée. Les exemples traités faisant apparaître quelques milliers de variables ont démontré l'efficacité de ces techniques pour les

problèmes non linéaires de très grande dimension.

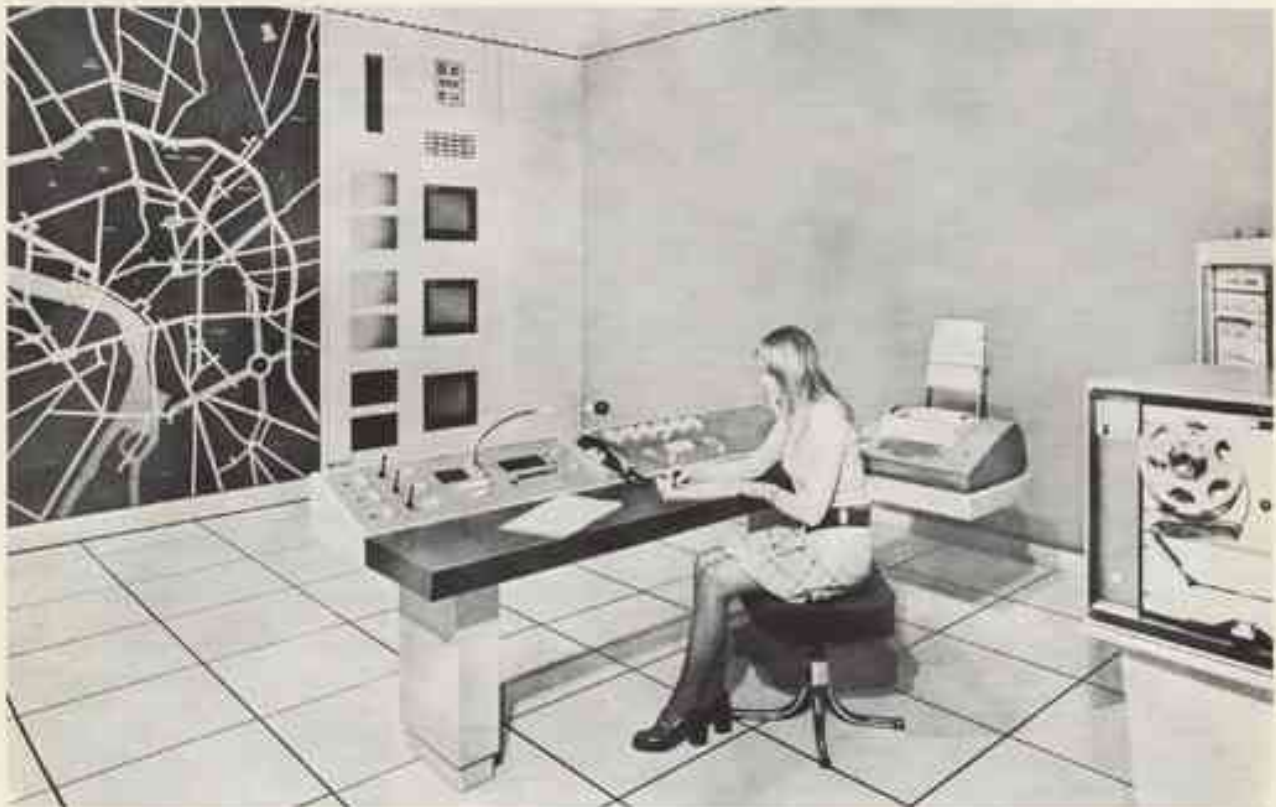
Dans le domaine de la commande "en ligne" des processus, on prendra l'exemple simple de la commande optimale d'un générateur de vapeur par une structure à deux niveaux. La pratique industrielle la plus courante consiste souvent encore à commander un processus composé de sous-systèmes par la juxtaposition de chaînes élémentaires de régulation sur chaque sous-système (fig. 5). Cette pratique subsiste même à côté de systèmes de commande plus modernes et prend le relais de ces derniers lors des démarrages et des arrêts ou lors d'un incident. Mais cette procédure nécessite un dédoublement des moyens de commande et présente donc des problèmes délicats de communication. Par contre, la décomposition du processus en sous-systèmes réglés séparément puis coordonnés de façon hiérarchisée conduit à une structure beaucoup plus séduisante puisqu'elle intègre des chaînes de régulation locales dans une commande globalement optimale. Elle permet avec la même structure de commander le système de façon classique, par des régulations localisées ignorant les interactions entre les différentes parties du processus, ou de façon globale en fournissant à ces "régulateurs locaux" une "commande de coordination" qui les oblige à harmoniser leur fonctionnement pour atteindre un but global.

Appliquée à la commande d'un générateur de vapeur cette approche conduit à faire apparaître une boucle de régulation de température et une boucle de régulation de pression dont les actions sont coordonnées par un niveau supérieur de commande (fig. 6).

Mais une telle structure autorise aussi de grandes simplifications dans la résolution mathématique du problème donc dans la synthèse des boucles de régulation locales, par rapport au problème global, même dans le cas où l'on suppose le système linéaire et où l'on choisit un critère quadratique, ce qui a été fait ici. Finalement, cette structure de commande hiérarchisée permet un fonctionnement satisfaisant en régulations locales (fig. 7) qui est hautement amélioré lorsque la coordination entre en jeu (fig. 8).

Les résultats obtenus sur les exemples évoqués démontrent déjà l'efficacité de cette nouvelle technique de commande. Lorsqu'on aura résolu un certain nombre de problèmes théoriques, elle deviendra un des outils principaux par lequel l'Automatique moderne pourra accroître son impact fructueux dans les secteurs industriels et socio-économiques, pour la commande des grands systèmes de production, de transport et de distribution.

conduite automatique du trafic urbain



Poste de commande centralisé pour le contrôle et la régulation automatique des feux de circulation de la ville de Toulouse, commandés par calculateur. (Photo Yan).

En matière de recherches relatives aux transports, l'attention se porte généralement sur les actions consacrées aux nouveaux types de véhicules. Cela pose cependant un certain nombre de problèmes complexes : étude et implantation de systèmes nouveaux, diversification et complémentarisation, extension des systèmes existants. Ces solutions se heurtent à de grandes difficultés techniques et économiques de mise en œuvre.

Un effort important est fait actuellement pour essayer de tirer profit au maximum des infrastructures existantes tout en améliorant l'exploitation des systèmes et ce par l'utilisation des techniques de l'automatique. Celles-ci, par leurs méthodes d'analyse et de synthèse permettent de conduire efficacement les études préliminaires de conception d'un système d'exploitation. Cependant la spécificité des problèmes rencontrés exige une adaptation et une évolution de ces méthodes, et fait sou-

vent l'objet d'une recherche théorique en automatique.

Le Laboratoire d'Automatique et d'Analyse des Systèmes s'est intéressé à de telles études. Son activité s'est tout d'abord centrée sur la conduite du trafic urbain, puis a été suivie par la réalisation d'un système de Conduite automatique du Trafic Urbain de la ville de Toulouse (système CONTRUT) sous la responsabilité d'une compagnie industrielle (photo 1). D'autres organismes ont abordé des études de même nature.

A l'étranger, elles sont activement poursuivies, notamment au "Road Research Laboratory" du Royaume Uni et de nombreuses réalisations de conduite du trafic faisant appel à des calculateurs fonctionnent.

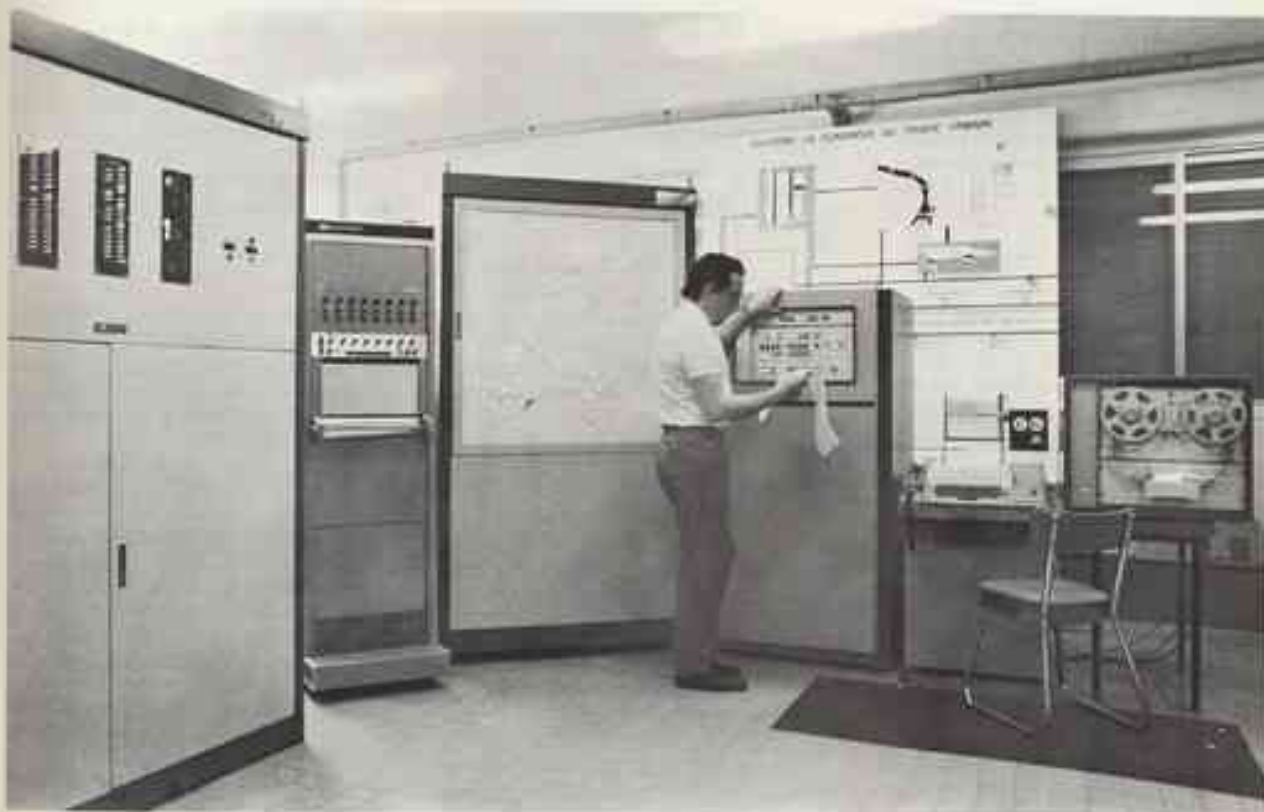
- L.A.A.S. - Toulouse.
- Laboratoire d'Automatique de l'ENSCM, Nantes.
- Centre d'Informatique de l'Université de Toulouse.

Quelques aspects du trafic urbain

Le trafic urbain se caractérise par la circulation, sur une infrastructure donnée, d'autobus et de véhicules qu'on ne peut commander directement, car chaque conducteur est considéré comme un organe de décision. Ce trafic est donc la résultante de conflits mutuels permanents. Plusieurs voies complémentaires sont possibles pour assurer l'organisation optimale de ces conflits :

- la réglementation,
- le plan de circulation qui agit par l'aménagement de l'espace disponible en interdisant certains itinéraires,
- les feux de carrefours qui assurent le partage du temps pour l'utilisation des zones conflictuelles.

La plateforme de recherche mise en œuvre pour le système CONTRUT s'est limitée au troisième type d'intervention : comment assurer, sur un réseau urbain, la gestion des feux de carrefours pour



Le simulateur de trafic (simulateur, visualisateur, calculateur et périphérique).

que l'écoulement global du trafic se fasse avec une perte de temps minimale.

Lorsque le trafic réel présent peut être écoulé, la commande consiste à gérer au mieux les conflits aux carrefours par adaptation de la répartition des durées de phases vertes et par adaptation de la synchronisation des débuts de phase verte. Lorsque le réseau est saturé en certains points, la commande vise avant tout à gérer le stockage des véhicules en excès pour maintenir le trafic à son débit maximum. Ce deuxième mode de commande du trafic urbain, en raison de ses difficultés extrêmes d'étude et de mise en œuvre, n'a pas constitué l'objet immédiat des travaux.

Evolution de la commande des feux de carrefour

Implantés pour des raisons de sécurité, les premiers feux de carrefour agissaient sur un trafic très modéré pour lequel

les temps d'attente n'étaient pas jugés significatifs. Avec l'augmentation du trafic, il a fallu se préoccuper de la qualité et de l'amélioration du service rendu aux usagers.

- Les premiers feux fonctionnaient avec des durées de vert indépendantes du trafic et déterminées empiriquement. Par la suite, des durées furent calculées "hors ligne" pour minimiser les temps d'attente des véhicules sur les bras du carrefour en supposant que le débit des véhicules était aléatoire à moyenne constante.

- Lorsque plusieurs carrefours sont rapprochés, l'arrivée des véhicules en certains bras présente un aspect pulsé. La recherche d'un réglage des feux améliorant la qualité de l'écoulement a conduit à leur synchronisation. Sur le site, les paramètres de réglage conservaient des valeurs constantes. Des méthodes d'optimisation (Retard-décalage,

TRANSYT, etc...) ont été développées pour aider l'ingénieur de trafic à définir la politique fixe qui minimise les temps d'attente. La commande des feux à politique fixe est une commande en boucle ouverte. Toute modification de cette commande est le résultat d'une remise en question, par l'ingénieur de trafic, des données qui ont servi à la détermination de la politique fixe précédente et dont les effets sont tardifs en regard de l'évolution du trafic. On sait en effet que celui-ci varie à chaque tranche horaire, à chaque minute même si l'on s'intéresse aux variables microscopiques de ce trafic.

- Aussi a-t-on implanté des dispositifs de commande en boucle "faiblement" fermée où la politique fixe appliquée sur le site est déduite d'une liste préalablement établie en fonction de l'heure courante ou en fonction des indications fournies par des capteurs de trafic implantés.

• Pour des carrefours isolés enfin, la commande des feux a été réalisée en boucle "plus fortement" fermée : les véhicules commandent directement la commutation des feux suivant une loi générale préalablement établie (commande adaptative). L'ingénieur de trafic ne se préoccupe alors, pour l'essentiel, que de l'établissement de la loi.

• La commande adaptative élargie à un ensemble de carrefours constitue une des solutions les plus performantes pour la commande en temps réel du trafic urbain. Toutefois elle a été peu développée jusqu'à présent en raison du coût de sa mise en œuvre à l'aide d'un matériel spécialisé à logique câblée. L'arrivée sur le marché de petits calculateurs industriels compétitifs avec ce matériel mais aux possibilités accrues et la tendance générale à l'accroissement du nombre de capteurs de trafic ont permis d'envisager l'application de nouvelles méthodes de commande en temps réel.

Commande en temps réel du système CONTRUT

Le réseau de capteurs étant supposé dense (plus de 2 capteurs par carrefour en moyenne), il est possible de caractériser le trafic de manière assez précise, de suivre ses fluctuations d'un cycle à l'autre et d'adapter les durées de vert en conséquence.

Une telle adaptation est réalisée pour les carrefours isolés à commande adaptative. Dans le système CONTRUT, cette adaptation est "contenue" pour que chaque carrefour d'une zone fonctionne à cycle commun lentement variable et pour que la synchronisation soit préservée. Celle-ci est définie sur la base des données de trafic filtrées tous les 8 cycles par une méthode d'optimisation utilisant un modèle de trafic simplifié exploité en ligne.

La politique de commande ainsi établie à chaque cycle par le calculateur en temps réel de petite taille, est transmise à un ensemble coordinateur (logique câblée) chargé de délivrer les impulsions de commande de basculement des feux de carrefours.

Méthodologie de la recherche

La mise au point des algorithmes de commande décrits et le test d'efficacité a priori des procédures adoptées ne peuvent, pour de nombreuses raisons pratiques, s'effectuer sur le site lui-même. Aussi l'utilisation d'une méthode de simulation peut rendre de grands services. La modélisation du trafic a été conduite au préalable par un recueil d'informations (photographies aériennes prises à chaque seconde).

Ainsi les lois d'avancement et d'arrêt des véhicules ont pu être restituées.

La simulation est réalisée par un "calculateur hybride" de trafic urbain (simulateur hybride) développé spécialement suivant le modèle retenu (photo II). Ce simulateur est piloté par un petit calculateur (CII 10010) en temps réel qui assure la gestion des données et met en œuvre les algorithmes de commande étudiés. L'intérêt du simulateur réside dans sa facilité d'emploi et sa rapidité de simulation. Il peut ainsi être utilisé par des ingénieurs de trafic conduisant un projet d'aménagement de feux de carrefours.

De l'identification et de la modélisation précédente a pu être déduit un modèle simplifié de trafic qui peut être implanté sur un petit calculateur en temps réel. L'exploitation de ce modèle à l'aide d'une méthode de programmation non linéaire permet le calcul du réglage optimal des feux. Par l'étude critique des propriétés du critère à optimiser (temps d'attente minimum des usagers) on a mis au point un algorithme hiérarchisé d'optimisation particulièrement utile pour le type de processus traité.

Ces études ont ainsi permis d'aborder par ses aspects les plus pratiques la conduite d'un processus complexe et flou dont nombre d'autres exemplaires existent et pour lesquels des commandes de type hiérarchisé devront être développées.

Prolongation des recherches sur le trafic urbain

L'expérience acquise tant sur les aspects théoriques que pratiques, ces derniers aspects ayant été transmis par l'exploitant et le réalisateur de CONTRUT, ont incité à aborder l'étude d'un type de commande à cycle libre des carrefours d'un réseau.

Il est reconnu que ce mode de commande fournit l'espoir des plus grandes performances, mais que les problèmes qu'il pose dans sa phase de développement sont également très importants.

Dans le contexte restreint où elle s'est inscrite, cette recherche a débouché sur la réalisation du système CONTRUT pilotant les feux de 35 carrefours de Toulouse à partir d'un calculateur situé dans un poste central de la ville. Outre les avantages divers offerts au service assurant l'exploitation des feux, ce système a permis de réaliser, dans des conditions semblables de trafic non saturé, une amélioration des temps de parcours pouvant atteindre 10 % (dans le cas de la ville de Toulouse) par rapport à ceux obtenus lors de la commande des feux prévalant antérieurement.

Mais du point de vue de la recherche, l'intérêt d'un tel système expérimental réside surtout dans les possibilités qu'il offre, par sa souplesse, de préciser la connaissance macroscopique du trafic.

Si les méthodes de simulation sont un moyen de contraindre le processus étudié pour le mieux connaître, elles n'en sont pas moins insuffisantes si elles ne s'inscrivent pas dans un dialogue avec la réalité. Ceci nécessite la collaboration d'organismes complémentaires à un laboratoire des sciences pour l'ingénieur.

Par ailleurs, des actions de recherche telles que celles entreprises ne peuvent garder leur sens que si elles s'insèrent dans un contexte de plus large réflexion où n'est négligé aucun des aspects équivalents, parfois pluridisciplinaires, susceptibles d'améliorer la qualité du service rendu à la collectivité.

détection et prévention des pannes dans les circuits logiques

Elaborer et construire des objets im-
pitoyable à chaque instant de contrôler
qu'ils satisfont aux fonctions auxquelles
ils sont destinés. Cette proposition reste
vraie lorsqu'il s'agit de vérifier le bon
fonctionnement d'un circuit logique et
d'y déceler la présence éventuelle
d'anomalies.

Une panne ou un défaut peut être défini
comme tout changement non souhaité
dans le circuit qui modifie la fonction
logique qu'il doit réaliser : on parle
ainsi de défaut de fabrication ou de
panne de fonctionnement.

De nombreux chercheurs et ingénieurs
se sont penchés sur la prévention ou la
mise en évidence de pannes de fonc-
tionnement. Les premiers travaux ont
été publiés vers les années 1960. Ils ne
se sont intensivement développés qu'à
partir de 1965, et en Juin 1971 s'est tenu
le premier symposium, Fault Tolerant
Computing, qui a lieu maintenant tous
les ans. Les Américains sont très nom-
breux sur ce sujet. En France, une di-
zaine d'équipes qui totalisent plus de
quarante chercheurs (industriels et uni-
versitaires) y travaillent aussi très direc-
tement.

Les progrès de la technologie et notam-
ment de la technologie des circuits
intégrés ont considérablement activé
ces recherches. Ces progrès se mani-
festent à plusieurs niveaux :

- La miniaturisation permet de réaliser
des circuits logiques qui comportent
des milliers de transistors, pour le même
encombrement que celui nécessitant autre-
fois par le boîtier d'un seul transistor ;
- la fiabilité a été considérablement
augmentée : le temps moyen entre
deux pannes est passé de quelques
dizaines de minutes à plusieurs milliers
d'heures ;
- la complexité des ensembles électro-
niques que l'on peut concevoir s'est
accrue dans les mêmes proportions ;
- la vitesse de fonctionnement des cir-
cuits a augmenté de plusieurs ordres
de grandeur : les temps de commutation
des circuits se sont réduits de la micro-
seconde à la nano-seconde.

Les méthodes automatiques

Cet accroissement en taille et en com-
plexité des circuits digitaux actuels
impose la mise au point de méthodes
automatiques et rapides de vérification
de fonctionnement qui seront utilisées
pour générer les tests à appliquer au
circuit et pour la mise en œuvre de ces

tests. Ces derniers seront les épreuves
appliquées au circuit pour connaître et
vérifier son comportement.

L'ensemble de détection est un ensem-
ble de tests qui permet de répondre à
la question : "le circuit est-il en panne ?".
L'ensemble de localisation est un en-
semble de tests qui permet de situer
l'emplacement de la panne. On conçoit
que le nombre de tests nécessaires à la
détection est moins important que celui
nécessaire au diagnostic.

On distingue deux types de tests selon
leur nature :

- les tests électriques sont statiques
(mesure de courant ou de tension conti-
nue) ou dynamiques (mesure de temps
de réponse) ;

- les tests fonctionnels qui font l'objet
essentiel de l'article permettant de
s'assurer que le circuit réalisé répond
aux conditions logiques souhaitées.

En ce qui concerne la nature des pannes
qui affectent un circuit logique, elles
peuvent être intermittentes ou perma-
nentes, logiques ou analogiques, uni-
ques ou multiples (c'est-à-dire peuvent
concerner une ou plusieurs connexions).
La plupart des résultats acquis ont trait
aux pannes logiques permanentes, com-
portant souvent une hypothèse supplé-
mentaire sur le type de panne : le défaut
se manifeste par un collage franc d'une
connexion, au niveau haut de tension
(collage à 1) ou au niveau bas (collage
à zéro), ce qui élimine d'autres types
de pannes tel le court-circuit entre deux
connexions.

Une deuxième classification apparaît
selon qu'il s'agit de logique combina-
toire ou séquentielle. Dans le premier
type, les valeurs de sortie du circuit ne
dépendent que des valeurs aux entrées
à l'instant considéré, alors que dans le
second, les valeurs de sortie dépendent
aussi de l'état interne de la machine.
Dans le cas d'un circuit combinatoire,
la première approche consiste à utiliser
une table de pannes qui fournit l'ensem-
ble des pannes correspondant à l'en-
semble des tests. Pour chaque confi-
guration d'entrée appliquée au circuit,
on recherche les pannes qui sont mises
en évidence. Si le circuit a n entrées,
la table de panne comporte 2^n lignes
et elle comporte 2^q colonnes si le cir-
cuit comporte q connexions qui peuvent
être collées à zéro ou à un.

- L.A.S. - Toulouse.

- Institut de Mathématiques appliquées de
Grenoble.

- Laboratoire d'Automatique de Mompel-
lier.

- Laboratoire d'Automatique de Grenoble.

Cette méthode devient inapplicable
pour des circuits de taille importante.
Pour réduire alors le nombre de pannes
à considérer, on étudie les classes d'équi-
valences entre des pannes. Considérons
par exemple le cas de la porte ET dont
la sortie vaut 1 si les deux entrées a
et b valent 1. Le collage de la sortie à
zéro ou d'une des connexions d'entrée
à zéro aura le même effet : les deux
pannes sont donc équivalentes. Mais
on peut aussi pour réduire le nombre
de tests rechercher, par des méthodes
d'analyse booléenne, les configurations
d'entrée qui permettent de mettre en
évidence une panne. D'une façon gé-
nérale, deux contraintes sont à res-
pecter : d'une part la "commandabilité"
du circuit (il faut pouvoir afficher aux
entrées d'une porte élémentaire les va-
leurs telles que l'effet d'une panne se
manifeste à la sortie de la porte) ; d'autre
part l'observabilité du circuit (il faut que
l'effet de la panne se propage jusqu'à une
sortie du circuit où il puisse être obser-
vé). La méthode dite du D algorithme
répond à ces deux contraintes. De plus,
elle ne nécessite pas que le circuit à tester
soit représenté par des équations boo-
léennes, ce qui permet d'éviter le pro-
blème du nombre d'informations à traiter.
Dans cette méthode, à chaque porte élé-
mentaire on associe des ensembles de va-
leurs entrées-sorties qui peuvent être
zéro, un ou D. La valeur D signifie que
lorsque le fonctionnement est normal,
la connexion est à un et, en présence
d'un défaut, cette valeur est zéro.

Dans le cas de circuits séquentiels, il se
présente d'autres difficultés car il faut
que le circuit soit dans un état interne
initial connu ; de plus, pour observer
les conséquences d'une panne, il faut
appliquer une séquence d'entrée. Les
approches utilisées pour les circuits
combinatoires ne peuvent pas s'appli-
quer directement. Une voie consiste
à transformer le circuit séquentiel en un
circuit combinatoire itératif constitué
par une cascade de circuits identiques
dont les entrées seront les entrées pri-
maires, accessibles directement à la
commande et les variables internes,
inaccessibles.

En outre, après la génération d'une
séquence fonctionnelle de test, il faut
s'assurer que cette séquence répond aux
contraintes du fonctionnement dyna-
mique.

Approche pratique du problème

Dans le cas de grands systèmes logiques, la mise en œuvre de ces méthodes d'analyse et de simulation reste difficile, voire impraticable, si l'on souhaite une détection des pannes à cent pour cent. Il faut donc considérer les problèmes de maintenance au niveau de la conception même du circuit. Dans ce but, on cherche à accroître la commandabilité et l'observabilité du circuit. La première démarche consiste à essayer de définir une structure du système qui suive un découpage en sous-ensembles fonctionnels, chacun de ces sous-ensembles étant accessible de la façon la plus autonome possible. Des sorties supplémentaires (points de test) sont ajoutées, dont l'observation permet la réduction des séquences de tests à appliquer au circuit. Mais on peut aussi insérer dans le circuit des modules logiques qui permettent non seulement l'observation mais le forçage du circuit dans un état interne donné. On aboutit ainsi à la conception de circuits faciles à tester ou à diagnostiquer. Cette propriété supplémentaire a pour conséquence un accroissement de l'ordre de quinze pour cent du matériel nécessaire.

Mais on peut songer aussi à tester une machine en cours de fonctionnement normal, c'est "le test en ligne". Il suffit de concevoir une logique de contrôle pour qu'une panne simple déclenche un signal d'alarme ou pour qu'elle se traduise (par l'emploi de techniques de

codage) par l'arrêt de la machine dans un état déterminé à l'avance. Dans tous les cas, ces améliorations exigent un matériel plus important.

Dès lors qu'une panne a été détectée par des moyens automatiques, il faut prévenir ses effets par des moyens eux aussi automatiques et réaliser des machines à autodépannage. La solution la plus directe est la duplication, la triplification même du circuit. On fait appel aux techniques de redondance, analogues au double circuit de freinage des voitures. Cette redondance peut prendre deux aspects : le premier, incompatible avec la détection de pannes, consiste précisément à masquer son effet ; le second consiste à disposer d'organes de secours auxquels il sera fait appel (la roue de secours en est un exemple). A cette redondance du matériel peut s'ajouter une redondance dans le temps : on effectue deux fois la même opération par un dispositif donné et on compare les résultats. Ces études doivent être menées par différents spécialistes car les domaines concernés sont divers : l'électronique intervient par exemple au niveau des dispositifs de commutation qui permettent une partie défectueuse par une partie saine équivalente ; la logique et la théorie des graphes sont utiles pour concevoir les chemins qui contournent les éléments en panne ; l'informatique est bien sûr concernée pour la définition de structures de machines, ainsi que la théorie des probabilités pour estimer la fiabilité.

Le problème essentiel pour toutes ces études est celui de la détermination du niveau auquel doit intervenir d'une part le diagnostic d'autre part la réparation.

On peut estimer à soixante pour cent le matériel supplémentaire nécessaire pour obtenir une dégradation douce, c'est-à-dire pour que le système bien que comportant des éléments défectueux, assure encore des tâches prioritaires. Cette dégradation douce est bien sûr adaptée aux missions que l'on attend du système.

Actuellement, de nombreuses équipes industrielles disposent de programmes de détection de pannes. Toutefois les limitations et les problèmes rencontrés lors de l'utilisation de ces programmes motivent largement les recherches en cours. Cet approfondissement des connaissances sur le test et le fonctionnement des systèmes digitaux doit conduire à :

- repousser les limites que présentent les méthodes actuelles ;
- définir des structures plus faciles à tester (commandabilité et observabilité) ;
- concevoir des machines tolérant des pannes (redondance, reconfiguration). Les applications de ces études concernent d'une façon générale tous les systèmes digitaux. Elles ont une importance primordiale dans certaines applications particulières : missions spatiales et commandes de processus en temps réel notamment.